

第一章 质数的建造、分布及检验

§1 前 言

你看过《数学传播》1984年12月号（32期）上第三十四届国际科技展数学得奖作品简介吗？这类展览都以是否“创新”为作品的评价，而在数学部门中，在《数学传播》所介绍的八个得奖作品中，就有五个属于“数论”的领域，占了得奖作品的一半以上，为什么？因为数论有它特别迷人的地方——那就是在极简单的规则中作极复杂的变化，一个类似的例子是围棋。围棋的规则极简单而其变化极为复杂，可以说是最迷人的一种棋类。相反的，海陆空军棋其规则极繁而变化又小，几乎没有人下了。又研究整数论所遭遇的问题及研究对象要比其它各门数学来得简单明了，不外乎是讨论正整数 $1, 2, 3, 4, \dots$ 的种种特殊性质，其变化之大，一直困惑许许多多的数学家。

以往整数论曾一直被人认为是最纯的纯数学，没想到由于最近美国有人利用找一个大量目的质数因子的困难性及其它质

VISIT...

LANZAROTE
Caliente.COM

数的一些特有性质,而设计了一种可公开传递(即不怕被敌方截获)且保密性极高的密码,引起了军方、工商界的莫大兴趣,有关这方面研究经费大为增加,我们在此简略地说明一下该密码的原理(对数论不熟的读者可参阅[1]或先阅下章)。在收报方甲先找出两个大的质数 p 、 q ,使 $m=pq$,及取任一与 $\varphi(m)=(p-1)(q-1)$ 互质的整数 a ,将此两数值 m 、 a 公开传递给发报方乙(甚至登报声明!),现假设发报方乙要将一信号(整数形式)拍给甲。设该代号为整数值 x (这是保密的且比 p 、 q 小很多),当然乙不能迳自公开拍发 x 给甲,而公开拍明码 $c \equiv x^a \pmod{m}$ 给甲方,现甲方收到明码整数 c 后设法译回到 x ,如何做到此一译码的工作呢?因甲方有资料 a 及 $\varphi(m)$, φ 为Euler函数, $\varphi(m)$ 表所有小于 m 且与 m 互质的正整数的数目。由假设 $a, \varphi(m)$ 互质,故有正整数 d 及负整数 b 使得 $ad + \varphi(m)b = 1$,这时甲方将收到的明码 c 作变换: $y = c^d \pmod{m}$;注意取 y 为小于 m 的整数,及由于 x 小于 p 及 q ,故 x 必与 m 互质,因而由尤拉(Euler)定理我们有 $x^{\varphi(m)} \equiv 1 \pmod{m}$,因此

$$\begin{aligned} y &\equiv c^d \equiv (x^a)^d \equiv x^{ad} \equiv x^{1 - \varphi(m)b} \\ &\equiv x \cdot x^{-\varphi(m)b} \pmod{m} \equiv x \pmod{m} \end{aligned}$$

现若两正整数 x 、 y 皆小于 m 且模 m 同余,故只有 $x=y$,因而甲方就可把密码收到了。

欲解此密码势必要知 d ,但要知 d 非要知 $\varphi(m)(=(p-1) \cdot (q-1))$ 也即要知 p 及 q ,找不到 p 、 q 就无法得 $\varphi(m)$ 及 d ,而硬要从 c 得出 x 就似乎很困难了。目前分解一个整数 n 的因子仍停留在近似硬试的阶段,由后面的“筛法”原理我们知道要从2, 3, 5, 7, ...一直试到小于 $\sqrt{n}$ 的质数为止,由[1]中可知若 n 为 ~ 50 位数(p 、 q 皆为25位数),则分解 n 要除 10^{25} 次,以每秒 10^6 次的电脑计算速度则将是一个 10^{11} 年的工作,若用特殊的快速法则来

进行也得要 10^{10} 次的运算，约4个小时电脑的计算时间，若 n 为一个100位的整数，则用目前最快速的电脑来操作运算也得要74年左右(中间还要保证机械没故障才行)，所以目前用这种方法来传递需保密的密码是相当安全的了。

质数可以说是整数的基础，由上面的应用我们欲充分利用质数必须要能建造很大的质数，要侦破上面密码的应用，我们也要知道质数的分布。本文就是针对此而需求作些浅显的介绍，希望能引起读者更大的兴趣，作更进一步的研究(而且可以保证的、是任何这方面的突破，无论在理论上与实用上都会引起强烈的反响)，光大我们祖先的光辉(因有的密码是利用“中国剩余定理”造成的)。

具体地讲，我们主要将介绍的是：1. 质数及质数表的制造；2. 如何构造任何一串列的大质数；3. 质数在整数中的分布或密度；4. 质数的检验。又本文的介绍之参考书〔1〕、〔2〕皆是新近出版的。

§ 2 质数表(筛法)及质数的制造

由于质数是不具有任何异于1及其本身的因子，所以早在纪元前200年左右古希腊学者Eratosthenes(以下简称爱氏)就为我们发明一个可找出所有质数的法则，称为筛法(Sieve method)。据说爱氏在找质数时，他把整数一一照序写在一片草质的纸上，凡是非质数他就用火在那位置烧一个洞，最后整片纸只留下密密麻麻的许多洞，很像一个筛子，故叫做“筛法”。它的原则如下：将正整数由1照大小依次排出一列或一矩阵(如图1是一正方形矩阵，其由1至100的整数组成)，则头一个

数为1，为非质数删掉，其下一个为2为质数留下，则2以后删掉所有2的倍数(4, 6, 8, 10, 12, 14, 16, 18...), 然后在所剩的数中大于2的第一个数3，其为一质数，继3以后删掉3的倍数(即6, 9, 15, 21...), 3之后5为第一个未被删掉的数其为一质数，删掉所有5的倍数(10, ...), 如此泡制，所留下的就是所有的质数(若只取1至 N 来筛，则如此所得的是1与 N 间所有的质数。在没有一张质数表时，我们怎样确定一给定的正整数 N 是否为质数呢?在 N 很大时若是以2到 $N-1$ 所有的质数一一来试除 N 会是件很耗时的事的，好在我们只需用2到 $\sqrt{N}$ 间所有的质数来试即可，这样一来就省了许多的除法。这是因为若 N 为非质数，则 $N=n_1 \times n_2$, n_1 及 n_2 为两个大于1的正整数，且必有一个数 n_1 或 n_2 不大于 $\sqrt{N}$ ，利用此一观察及筛法我们很容易把所有不大于某个正整数 N (N 不是很大时)的质数找出来。例如我们要列出100以下所有的质数，我们只需用小于 $\sqrt{100}=10$ 以下的质数2, 3, 5, 7用筛法把所有小于100的质数找出如下：

2, 3, 5, 7, 11, 13, 17, 19, 23,
29, 31, 37, 41, 43, 47, 53, 59,
61, 67, 71, 73, 79, 83, 89, 97.

目前我们有10亿($=10^9$)以内的质数表。而利用大质数的密码为了防止敌人利用电脑来侦破，往往是用有50位的质数，所以用筛法来建造大质数仍无法合于实用。

又用筛法制造大的质数时的一个缺点是每次要列出或利用许多的正整数来求得，这在实用上很不方便。我们很想随意制造一些很大的质数，该如何做呢？最理想的是我们能找出一个公式或一个式子，只要把适当的数代入就可得出很大的质数，我们中国人很早有一个有关质数的制造及拣选，认为一个正整数 n 为质数的充要(充分及必要)条件是 n 可整除 2^n-2 ，即

$n|(2^n - 2)$, 例如 $n=5$, $2^5 - 2 = 30$, $5|30$, 这个结论当时并没正式证明, 而事实上现今可证明此条件 $n|(2^n - 2)$ 为必要的(这可

1	2	3	4	5	6	7	8	9	10
11	12	13	14	15	16	17	18	19	20
21	22	23	24	25	26	27	28	29	30
31	32	33	34	35	36	37	38	39	40
41	42	43	44	45	46	47	48	49	50
51	52	53	54	55	56	57	58	59	60
61	62	63	64	65	66	67	68	69	70
71	72	73	74	75	76	77	78	79	80
81	82	83	84	85	86	87	88	89	90
91	92	93	94	95	96	97	98	99	100

图1 1至100的筛法求质数

由所谓的小费马 (Fermat) 定理来证明, 它是说: 若 a 与 k 为两无公因子的正整数, 则 $a^k \equiv a \pmod k$, 因此 n 为奇质数时此条件为必要的立即可见, 而当 $n=2$ 时则 $n|2^n - 2$ 明显成立)。但此条件却非充分的, 不过至少以当时的计算能力来检验此一真实性, 恐怕也不是件很容易的工作, 因为事实上对比小于341的正整数 n , 若 n 不是质数则 $n|2^n - 2$, 例如 $n=15$, $2^{15} - 2 = 32766$, 此数不可能被15整除, 但 $n=341 (= 11 \times 31)$ 时, 用现代的任何小型电子计算机 (Programmable Pocket Calculator) 可得知341除 $2^{341} - 2$ 的剩余为0 (注意, 对于这样一个具有101位数的商, 就不是可由此一类机子可算出了, 这也难怪我们的老祖先在只有算盘代劳的劣势下, 更是心有余而力不足了)。

但很早就有人证明具有 $4n + 1$ 形式的质数是无穷多的 (读者

不妨试证一下！提示：任何一个整数 m 都可表成 $4n$, $4n \pm 1$, $4n + 2$ 的形式中之一），又对 $f(n) = n^2 + n + 41$ 对此一二次式而言，在 $n = 0, 1, 2, \dots, 40$ 连续41个整数的值时皆为质数，但很明显 $n = 41$ 时， $f(n)$ 有一因子为41，或许有读者会问可不可能找到一个多项式 $f(x) = a_0 x^k + a_1 x^{k-1} + \dots + a_k$ ，系数 $a_0, a_1, \dots$ （皆为整数）使得对任何的正整数 n , $f(n)$ 皆为质数？这个问题的答案不难知是否定的。其证如下：设 $f(n_0) = P$ 为一质数，则对任何整数 m , $f(n_0 + pm) - f(n_0) = \sum_{i=0}^k a_i [(n_0 + pm)^i - n_0^i] \equiv 0 \pmod{p}$ ，又方程式 $f(n_0 + pm) = 0$ 及 $f(n_0 + pm) = \pm p$ 至多有 $3k$ 个整数解 m ，故当 m 充分大时 $f(n_0 + pm) - f(n_0) \neq 0, \neq \pm p$ ，且为 p 的倍数，故为非质数。

在上面我们证明了没有一个多项式 $p(x) = a_0 x^n + a_1 x^{n-1} + \dots + a_n$ 存在使得对所有的整数 m , $f(m)$ 永远取质数值，但另一方面我们也不难证明对 $p(x) = 6x + 1$ ，必有无穷多的质数 $p(m)$ （证明参看〔1〕）。在这方面一个艰深的问题是一个二次或以上的整系数多项式，即 $p(x) = a_0 x^n + a_1 x^{n-1} + \dots + a_n$, $n \geq 2, a_0 \neq 0$, $a_i (i = 0, 1, 2, \dots, n)$ 皆为整数，能否证明 $\{p(m)\}_{m=1}^{\infty}$ 必有无穷多个质数？如果读者中有人能解决此一问题，他的大名一定会记在数学史上。

所以我们知道若要建造一个函数 $f(x)$ 使得 $f(n)$ 在 n 为正整数时，总为质数，想由多项式的形式中去寻求，是徒费其劳的。下面我们就介绍密尔史氏(W. H. Mills)的一个建造法(1947年)，其基本构想是基于1937年英格翰氏(A. E. Ingham)的一定理：若把所有的质数以大小渐增地排列并以序列 $p_1, p_2, \dots, p_n, p_{n+1}, \dots$ 或 $\{p_n\}_{n=1}^{\infty}$ 表之，即 p_n 表示第 n 个质数，则不等式： $p_{n+1} - p_n < k p_n^{\frac{1}{k}}$ 成立，此处 k 为一定正整数，依据此定理密尔史氏先证明了

下面的一个结果:

引理1.1 设 N 为一大于 k^8 的正整数, 则 N^3 与 $(N+1)^3-1$ 之间必有一质数.

证明: 设 p_n 为小于 N^3 的最大质数, 则

$$\begin{aligned} N^3 < p_{n+1} < p_n + kp_n^{\frac{5}{8}} < N^3 + kN^{\frac{5}{8}} \\ < N^3 + N^2 < (N+1)^3 - 1 \end{aligned}$$

令 q_0 为一大于 k^8 的质数, 则依据上面的引理, 我们可找到一无穷多的质数列 $q_0, q_1, \dots$,

$$\text{其满足 } q_n^3 < p_{n+1} < (q_n+1)^3 - 1 \quad (1)$$

$$\text{令 } u_n = q_n^{3^{-n}}, v_n = (q_n+1)^{3^{-n}} \quad (2)$$

则由(1)及(2)可得

$$v_n > u_n, u_{n+1} > q_{n+1}^{3^{-(n+1)}} > q_n^{3^{-n}} = u_n \quad (3)$$

$$\begin{aligned} \text{及 } v_{n+1} &= (q_{n+1}+1)^{3^{-(n+1)}} < (q_n+1)^{3^{-n}} \\ &= v_n \end{aligned} \quad (4)$$

因此 $\{u_{n+1}\}$ 为一渐增的数列, 且其为有界的, 因而极限

$$\lim_{n \rightarrow \infty} u_n = A \quad (5)$$

存在.

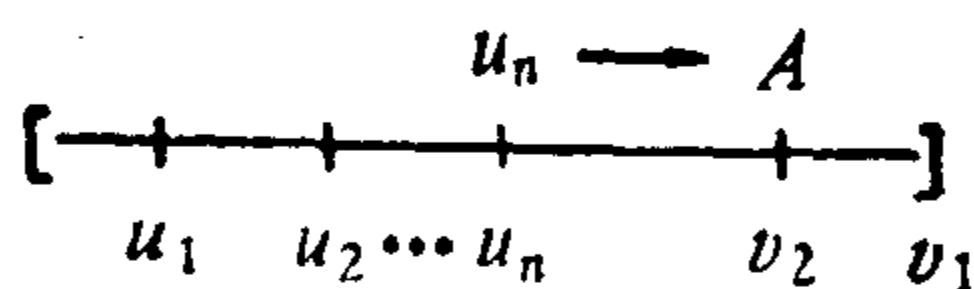


图2

接着密尔史氏证明了下面的定理, 为此先介绍一个定义及一记号.

定义1.1 称一个函数 $f(x)$ 为质数生成表示函数, 若对任何正整数 n , $f(n)$ 为质数.

$[R]$ 表示为一不大于 R 的最大整数, 例如 $[3.15] = 3$.

定理1.1 设 A 为引理1中所定, 则 $f(x)=[A^{3^x}]$ 为一质数生成表示函数.

证明: 由式(3), (4)及(5)可得知

$$u_n < A < v_n$$

$$\text{或 } q_n < A^{3^n} < q_{n+1}$$

因而依记号 $[\quad]$ 的意义, 立即得

$$[A^{3^n}] = q_n$$

于是我们证得 $f(n)=[A^{3^n}]$ 皆为质数.

读完上面定理的结果, 有些读者或许会看出它的一个缺陷, 就是这个制造大质数的函数是理论性的一个表示而已, 真正的值仍是无法得出, 因为我们不知 A 到底是多大? 只知道它是 $\lim_{n \rightarrow \infty} u_n$ 的极限值(存在且为有限的), 所以要制造大质数仍只有靠硬来的筛法了. 但有些具特殊形式的数中, 较容易得出一些大的质数(这是下节要谈及的).

我们知当 n 充分大时, n^3 与 $(n+1)^3-1$ 两数之间必有一质数存在, 和此现象似乎相冲突的是对任何给定的大的数 l 我们可以建造一个区间, 其长变为 l , 且在其间无一整数为质数. 比如我们想找一百万个连续的正整数, 其皆为合数(非质数及1的整数), 取 $q(n)=(10^6+1)!+n$, 则当 $n=2, 3, 4, \dots, 1000001$ 时为合数, 读者不妨试一下证明此一般性的结果, 由前面所提及的英格翰氏定理知一个无质数的区间的大小与两接连的质数间隔大小有关. 值得一提的是由它可推出有名的 Bertrand 臆测: 即在任何两整数 n 及 $2n$ 之间必有一质数, 即

$$p_{k+1} < 2p_k$$

另一有关质数分布的性质是在区间 $[n, 2n]$ 内的质数数目, 和 $[1, n]$ 间质数数目的大小是同级的 (order), 此一事实可由质

数分布的渐近函数来说明，今以 $\pi(x)$ 表示小于 x 的质数数目，则

$$\pi(x) \sim \frac{x}{\log x}$$

$$\begin{aligned} \text{故 } \pi(2x) - \pi(x) &\approx \frac{2x}{\log x + \log 2} - \frac{x}{\log x} \\ &= \frac{2x}{\log x} \left(1 + \frac{\log 2}{\log x}\right)^{-1} - \frac{x}{\log x} \\ &\approx \frac{2x}{\log x} \left[1 - \frac{\log 2}{\log x}\right] - \frac{x}{\log x} \\ &= \frac{x}{\log x} - \frac{2x \log 2}{\log^2 x} \\ \frac{\pi(2x) - \pi(x)}{\pi(x)} &\longrightarrow 1 \quad (x \longrightarrow \infty) \end{aligned}$$

§3 质数的分布

上节中对质数间的间隔或数目多少我们用到质数数目函数 $\pi(x) \sim \frac{x}{\log x}$ 一事实。我们现对质数此一重要性质作一些讨论。今任一大于1的正整数在整个整数集的分布情况，如它不是质数就是合数，好像没什么机率可言，严格讲我们将可由质数数目的渐近函数 $\pi(x)$ ，得知对任一正整数其为质数的概率为0，从我们熟悉的一事实开始：正整数列中每第2个数可被2整除，每第3个数可为3整除等等。如何利用此简单明了的现象去求得近似 $\pi(x) \sim \frac{x}{\log x}$ 的结果，是我们在下面要探讨的。这儿的讨论

并不严谨，但却由直觉的概率来说明，首先我们要求的是对一任意的正整数其可被质数 p_i 整除的概率为 $1/p_i$ ，这是因为由1开始每第 p_i 个数可被 p_i 整除，故一个数可被 p_i 整除之概率为 $1/p_i$ ，不被 p_i 整除之概率就为 $1 - 1/p_i$ 。我们不妨视被两个不同质数整除是两独立事件（不过对所有的质数视其皆为独立是不可能的，但可认为几乎是独立的；例如取一数可被2整除，不能帮助你去推测它是否能被3整除）。则对任一给定的整数 x 其不为任何小于其本身的质数整除的概率为

$$\begin{aligned}\omega(x) &\approx \left(1 - \frac{1}{2}\right)\left(1 - \frac{1}{3}\right)\cdots \\ &\approx \prod_{p_i < x} \left(1 - \frac{1}{p_i}\right)\end{aligned}\quad (6)$$

此时的 x 也都表示一质数。但依据筛法的原则，我们只该在上面乘积的式子中取小于 $\sqrt{x}$ 的质数 p_i ，但由于如此得出的乘积在分析问题计算中并没引起多大数值上的差别（ x 很大时），所以我们可不计较此改进了。如果对式(6)两边取对数可得

$$\log \omega(x) \approx \sum_{p_i < x} \log \left(1 - \frac{1}{p_i}\right)$$

这时如果我们同意依据 $\log(1+y) = y - \frac{y^2}{2} + \frac{y^3}{3} + \cdots$ 的展开式在 y 很小时取 $\log(1+y) \approx y$ 的近似值代入式(6)中，就可得下面在 x 很大时的近似表示

$$\log \omega(x) \approx \sum_{p_i < x} -\frac{1}{p_i}\quad (7)$$

我们仍想把此一表示化成一个更具体而简洁的形式。参照 $\omega(x)$ 的定义及试想上面式子(7)中项 $\frac{1}{p}$ ，其出现的概率为 $\omega(p)$ 。

此一想法使得我们可把式(7)改写为

$$\log \omega(x) = - \sum_{n=2}^x \frac{\omega(n)}{n} \quad (8)$$

将此式子表成积分形式可得

$$\log \omega(x) \approx - \int_2^x \frac{\omega(n)}{n} dn \quad (9)$$

为了去掉负号作变换 $A(x) = \frac{1}{\omega(x)}$ 则由上式可得

$$\log A(x) \approx \int_2^x \frac{dn}{n A(n)} \quad (10)$$

如果两边取微分, 得

$$\frac{A'(x)}{A(x)} \approx \frac{1}{xA(x)}$$

$$\text{即 } A'(x) \approx \frac{1}{x}$$

因而形式上得

$$A(x) \approx \log x$$

于是

$$\omega(x) \approx \frac{1}{\log x} \quad (11)$$

若视此为质数密度的平均值, 则对于小于或等于 x 的质数数目函数 $\pi(x)$, 就可用下面的近似式来表示了:

$$\pi(x) \approx \int_2^x \frac{dx}{\log x} \approx \frac{x}{\log x}$$

§4 有关质数检验的一些结果

要检验一个给定的正整数 n 是否为质数, 筛法是最直接的一

法子,但它要用所有小于 $\sqrt{n}$ 的质数来试除(大约有 $2\sqrt{n}/\log n$ 个质数),这对大的 n ,即使用电脑来计算也不切实际,我们知道要建造大的质数也不是件容易的事,但总有些有耐心及聪明的人在那探求找大质数,以下我们就介绍一些有关的尝试。

和费马(Fermat,提出有名的费马猜想的)及笛卡儿(Descartes)为朋友的法人墨森尼(Mersenne)曾列了一个具有形式 $M_n = 2^n - 1$,但为质数的一个表, M_n 就称为墨森尼数。很明显,若 n 为合成数,则 M_n 不可能为质数,所以只有 n 为质数时, M_n 才可能为质数,已知的墨森尼质数有

$$M_2 = 2^2 - 1 = 3, M_3 = 2^3 - 1 = 7,$$

$$M_5 = 2^5 - 1 = 31, M_7 = 2^7 - 1 = 127,$$

$$M_{13} = 2^{13} - 1 = 8191, M_{17} = 2^{17} - 1 = 131071,$$

$$M_{19} = 2^{19} - 1 = 524287, M_{31} = 2^{31} - 1 = 2147483647.$$

M_{31} 为质数此一事实是在1750年为尤拉所证明,且为1876年前所知的最大质数,而 $M_{11} = 2^{11} - 1 = 2047 = 23 \times 89$, $M_{29} = 2^{29} - 1 = 536870911 = 233 \times 1103 \times 2089$,所以不是对所有的质数 p , M_p 亦为质数。在1876年鲁克斯(E. Lucas)发现了一个法则间接地证明一个 M_p 数是否为质数,他主要应用了以 M_p 为模的整数环(ring),在此环中两个数的加减乘和往常一样,只是若结果的值不在0到 $M_p - 1$ 的范围内时,以用 M_p 除该数所得的为非负的剩余;例如以 M_5 为模的整数环中, $5 \times 7 \equiv 3 \pmod{M_5}$ 的一些基本性质。他的结果被近代的赖莫尔氏(D. H. Lehmer)精简为下面的法则:

若 p 为大于2的质数,定义 $L_1, L_2, \dots$ 如下:

$$L_1 = 4, \text{ 及对 } n \geq 2$$

$$L_n = L_{n-1}^2 - 2 \pmod{M_p}$$

则 M_p 为质数当且仅当

$$L_{p-1} \equiv 0 \pmod{M_p}$$

我们看如何用上面法则来检验 $M_{11} = 2^{11} - 1 \equiv 2047$ 是否为质数？因 $L_1 = 4, L_2 = 14, L_3 = 194, L_4 = 788, L_5 = 701, L_6 = 119, L_7 = 1877, L_8 = 240, L_9 = 282$ ，及 $L_{10} = 1736$ ，于是 $L_{10} \not\equiv 0 \pmod{2047}$ ，因而 M_{11} 不可能为质数。

在同年(1876)鲁克斯用上法检验 M_{127} ，他计算了 $L_1, L_2, \dots, L_{126}$ ，发现了 $M_{126} \equiv 0 \pmod{M_{127}}$ ，因而创下了打破保持了75年 M_{31} 为最大质数的纪录，自从1950年开始进入电脑的时代后，利用鲁克斯法则及二进位的计算原理，很多很大的质数陆续被得出。例如1952年证明了有687位数的 M_{2281} 为质数，1971年证明 M_{10937} 为质数其有6002位数，在1978年由两位18岁加州州立大学的学生共同证得具6533位的 M_{21701} 为质数。最近1983年有人证得具39751位数的 M_{132049} 为质数（在证明此一事实中，曾利用 $M_{132049} + 1$ 为2的幂次），如果读者中有人想在数论中留一席之地，不妨试试找出一个比此数更大的质数，这可能相当的困难。有关鲁克斯法则的证明及应用，我们在下章还会谈到。

在[1]中也提到一个有趣的事实：即 m 若不是质数，则至少有一半以上的数从2, 3, ... 到 $m-1$ 不能满足 $a^{m-1} \equiv 1 \pmod{m}$ ，这点说明若 m 为质数，则从2, 3, ..., $m-1$ 任取一数 a ，使得 $a^{m-1} \equiv 1 \pmod{m}$ 成立的机会相当大（至少是 $\frac{1}{2}$ ），但至今仍无一简便的法则找出一个与 m 互质的 a ，这是读者可以动手动脑试一试的问题。

我们提供下面有关质数的一些有趣性质，作为本文之结束。

§ 5 证明: $\sum_{i=1}^{\infty} \frac{1}{p_i} = \infty$

(即所有质数级数之和为发散的)

所有正整数之和: $1 + 2 + \cdots + n + \cdots$ 显然是无穷大的(即为一发散级数), 但如问所有正整数之倒数形式之级数 $\sum_{n=1}^{\infty} \frac{1}{n}$ 是否为发散? 解答此问题只要稍用些技巧就可达到, 但如可利用级数与积分之关系, 则也很容易证明此级数为发散的, 同样的问题对所有的质数(读者首先证明此集合为一无穷集的事实!), 我们也有下面类似的结果, 但它的证明就难得多了。它的证明也有几种, 但我们介绍的是一很初等的证明。

定理1.2 设 $p_1, p_2, \cdots p_k, \cdots$ 为所有质数且照递增秩序排列的数列, 则级数 $Q = \sum_{i=1}^{\infty} \frac{1}{p_i}$ 为发散的, 即 Q 的值为无穷大。

证明: 我们将用反证法, 即设 Q 为一收敛的级数并试由此导出矛盾。今假设 Q 收敛于一有限正数 s , 则由收敛值的定义, 我们可找到一个整数 k , 使得 Q 的 $k-1$ 项部分和

$$Q_{k-1} = \sum_{i=1}^{k-1} \frac{1}{p_i} < s - \frac{1}{2},$$

$$\text{而 } Q_k = Q_{k-1} + \frac{1}{p_k} > s - \frac{1}{2}$$

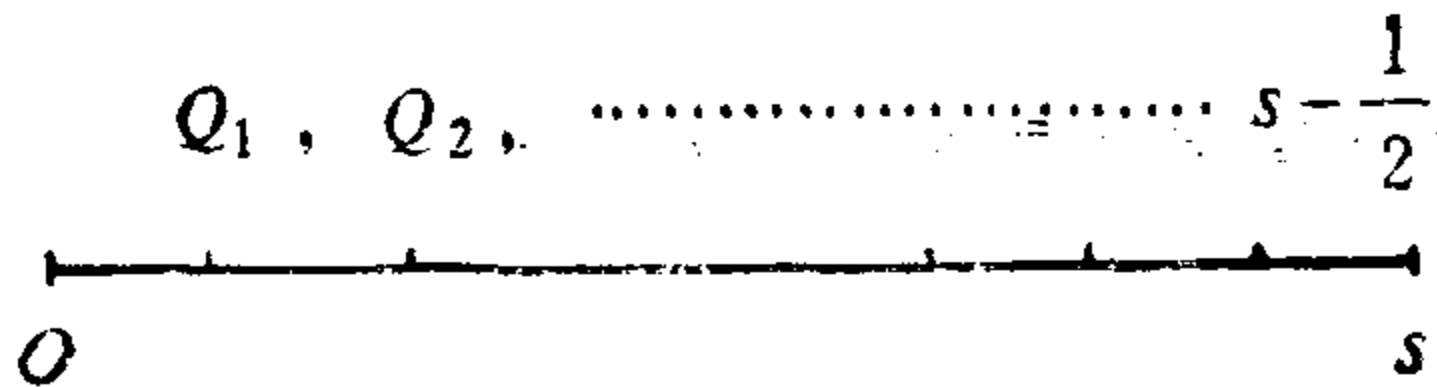


图3

(此处 $\frac{1}{2}$ 的选择并无特别的意义, 任何一个小的正数量都可)

于是有

$$p_k = Q - Q_k = \frac{1}{p_{k+1}} + \frac{1}{p_{k+2}} + \cdots < \frac{1}{2} \quad (12)$$

就此 k , 我们对任何一正整数 x , 定义 $N(k, x)$ 为所有不大于 x 的正整数, 其不被任何大于 p_k 的质数整除者的集合, 我们以 $|N(k, x)|$ 表集合 $N(k, x)$ 中元素的数目, 例如 $k=4$, 则所有大于 p_4 的质数集 $=\{p_5, p_6, p_7, \cdots\}=\{11, 13, 17, 19, \cdots\}=T_4$. 于是 $N(4, 10)=\{1, 2, 3, \cdots, 10\}$, 故 $|N(4, 10)|=10$, 而 $N(4, 15)=\{1, 2, 3, 4, 5, 6, 7, 8, 9, 10, 12, 14, 15\}$ 共13个元素(这是因为由1到15中除了11、13外没有一个正整数可能具有集合 T_4 中任一元素为因子者).

今对任一正整数 z , 我们都可先将它表成质因子的幂次的乘积:

$$z = p_1^{\alpha_1} p_2^{\alpha_2} \cdots p_i^{\alpha_i}$$

其中 $\alpha_1, \alpha_2, \cdots, \alpha_i$ 为整数. 若将所有具偶数幂次的项与将奇数幂次减1的项的乘积以 w^2 表之, 其余项的乘积以 v 表之, 则

$$z = w^2 v$$

其中 w, v 皆为整数, v 中的质因子的幂次皆为1. 如 $z=2^7 3^4 5^2 7^5$, 则 $z=(2^6 3^4 5^2 7^4) \cdot (2 \times 7) = (2^3 \cdot 3^2 \cdot 5 \cdot 7)^2 (2 \times 7)$. 现在 we 看什么样的正整数 y 可能属于 $N(k, x)$, 很显然它们必须满足:

(1) $y \leq x$.

(2) 所有 y 的质因子必来自 $\{p_1, p_2, \cdots, p_k\}$ 中.

依据上面两个条件, 我们先来估计一下 $|N(k, x)|$ 的大小. 由于对任何 $y \in N(k, x)$, 则有 $y = w^2 v \leq x$ 及 $v \geq 1$,

故 $w \leq \sqrt{x}$

而因 v 的因子来自 $\{p_1, p_2, \dots, p_k\}$ 且皆为1次幂, 故所有可能为 v 的数至多有 2^k 个, 综合以上的分析, 我们得知在符合条件(1)、(2)下, 至多有

$$[(\sqrt{x})] \times 2^k \text{ 个 } y$$

(其中 $[a]$ 表不大于 a 的最大整数)

可能属于 $N(k, x)$, 于是

$$|N(k, x)| \leq 2^k \sqrt{x} \quad (13)$$

故对于任一正整数 x 由1到 x 的正整数中, 其中 $|N(k, x)|$ 不可能被 $p_{k+1}, p_{k+2}, \dots$ 中任一个数整除, 剩下的共有 $x - |N(k, x)|$ 个, 每个数都可能具有 $p_{k+1}, p_{k+2}, \dots$ 中某些数为因子. 现从另一角度来估计 $x - |N(k, x)|$ 的大小, 在 $\{1, 2, \dots, x-1, x\}$ 中很明显至多有 x/p_{k+1} 个数可被 p_{k+1} 整除, x/p_{k+2} 个数可被 p_{k+2} 整除, $\dots$ 等等. 所以由此及(12)可得

$$\begin{aligned} x - |N(k, x)| &\leq \frac{x}{p_{k+1}} + \frac{x}{p_{k+2}} + \dots \\ &\leq x \left\{ \frac{1}{p_{k+1}} + \frac{1}{p_{k+2}} + \dots \right\} \\ &< \frac{x}{2} \end{aligned}$$

因而

$$|N(k, x)| > \frac{x}{2} \quad (14)$$

于是由(13)、(14)两式得

$$\frac{x}{2} < |N(k, x)| < 2^k \sqrt{x} \quad (15)$$

并注意此不等式的导至与 x 的大小无关, 特别取 $x \equiv 2^{2n+2}$ 时, 由(15)将得不等式:

$$2^{2n+1} < |N(k, 2^{2n+2})| < 2^{2n+1}$$

此为不可能。此一矛盾亦证明了 $Q = \sum_{i=1}^{\infty} \frac{1}{p_i}$ 必须为发散才行。

研究问题：如果我们限制质数的形式结果如何呢？此似乎是个新的问题，具体的我们可以问：设 $\{p_k\}$ 为具 $4n+1$ 形式的所

有质数序列，则 $\sum_{i=1}^{\infty} \frac{1}{p_i}$ 是发散抑或为收敛？

参 考 书 目

1. 《整数论及其应用》，杨重骏、杨照崑编著，东华书局，1983。
2. M. R. Schroeder, "Number theory in Science and Communication" Springer-Verlag, 1984.

第二章 数论在密码上的应用

§1 前言

数论，顾名思义，是一门研究数学性质的学问。一般所谓的数论，特指正整数（即自然数）的许多性质，像质数的分布、方程式的正整数解、韩信点兵及进位法都包括在数论里面。我们在小学时候学的分解因数、最大公约数也是数论的一部分，可惜因为数论在日常生活中没有什么直接的用处，在中学数学里很少提到数论，一般被认为是一种“纯数学”，深而无用。可是“无用之用，真乃大用”，终于在20世纪70年代后期，几个电机工程师用数论的一些基本定理，制成了一种新的密码。这种由数论所作成的密码与以前人们所用的密码，有着根本性质的不同，可说是密码史上一个空前的革新。

密码通讯在军事上的用途是大家都知道的，但由于交通的发达，在分秒必争的工商业社会里，商业上的情报也已成为商业盈余的主要依靠。比如说，有人早几个小时知道什么公司有了一种新的发明，或某两个公司计划合作；或某地区有物价的大波

以往人们在军事上所用的密码其基本的形式在于“代换”与“置换”。比如说，我要发出下面一个消息给你：

我就先把这几个字换成数字，即一般电码本上的代码，假定“我”字的代码是3314，“有”字的代码是1432，“一”字代码是0001等等，则上面那句话就成了

代换密码是把0, 1, 2, ..., 9十个数字互换, 比如我们可以把0换成2, 1换成3, 等等, 若用群论的符号表示, 上面的代换可写成

这个表示法是上行为0, 1, 2, ..., 9, 而下行是他们代换成的新数字, 即 $0 \rightarrow 2$, $1 \rightarrow 3$, $2 \rightarrow 5$, ...。因此刚才的电码若用G法代换, 则成了

这时一个不知道这个代换规则的人看到了上面的信号，他就不能从电码本子里找出它的原意了。置换法在于把密码排成一种双方都知道的形式，如下图



则发出的信号为755772433216623。同样地,不知道这种特定图案的人,很难解开原来的信息。

以代换法为例,像G这类的转换可以有 $10! = 3628800$ 种不同的变化,假定我们可以在一分钟内试一种代换,又假定我们的运气中等,在试到一半时即 $10!/2 = 1814400$ 时可以成功,则在不吃、不睡、不错的情形下,我们要试3年零165天,等迷解出来的时候,仗早已打完了。但这只是最基本的密码而已,在生死关头,更难解的密码必然出笼,例如在代换法中,若两位两位的代换(即 $00 \rightarrow 79, 01 \rightarrow 85, \dots$)则其变化可达 $100! = 9.32 \times 10^{57}$ 种之多,如果我们再用硬试的方法,则一百万人同心协力也得用 6×10^{48} 年才能试出谜底,可是地球的年龄不过只有 5×10^9 年而已。

因此解密码,都不用硬试的方法去解。一般可用统计的方法,根据名字(或字母)出现的频率及发生的事件加以分析,例如在英语中,各字母出现的频率按多少排列是

e, a, o, i, d, h, n, r, s, t,

u, y, c, f, g, l, m, w, b, \dots

因此一个出现次数最多的符号就很可能代表 *e*, 出现次多的符号就很可能代表 *a*, 并以此类推。但道高一尺,魔高一丈,如果不断地变化代换的方法,比如说前一百个字用代换法 G_1 , 第二个一百个字用代换法 G_2 , 则频率方法亦失去功效。但是魔高一丈恶魔更高一丈,重赏之下,天下尚没有绝对解不开的密码。二次大战时,美国密码兼统计学家弗立得门(W. F. Friedman),在一年半的时间完全破获了日军的密码。在中途岛一战,美国海军以劣势的军力,大胜日本皇家海军(读过参加该战役的日本军官渊田美津雄及奥官正武所著的“中途岛”一书吗?)。另外英国也几乎在同时解开了德军的密码,做到了知己知彼的优

勢地位（知己知彼并不一定百战百胜，但总比不知彼要好得多）。

过去这类密码的特质（也可以说是缺点）在于它们是关闭式的制解法，即收发双方都必须同时知道这种密码的构造。因此如果在一通讯系统中有一个联络站被间谍渗入或被敌人占领，则密码的机密可能全盘暴露。而现在用数论的密码却是公开式的（Public-Key Cryptography），即是只有收方知道密码的解法，发方只需要知道做法而已，而且这种制法可以公开。因此即使发方被捕，敌人仍榨不出解密码的机密来，其程序是这样的：

1. 收方先告诉发方如何去把情报做成密码（敌人也听到了这个做法）。

2. 发方依法发出情报的密码（敌人也听到了这个信号）。

3. 收方解开此密码为原信息（但敌人却解不开此密码）。

当然把收发互换，就可以互通信息了。刚才说过这种方法最大的好处就是只有一方知道解码的秘诀，比以前收发双方都知道秘诀的保密性高多了，自从这类的密码法发表之后，在美国军事界、教育界、工商业界引起了一个轩然大波。对大多数的数论家而言则是一则以喜，一则以惧。喜的是，皇天不负苦心人，一向不容易找到饭吃的数论家突然成为许多经费充裕的军工商界所争取的对象。费马、尤拉以及那些一生穷困而已作古了的数论家可以含笑九泉了。但惧的是，由于这些理论在军事通讯上的价值，有关这方面的新发现已被视为国防机密而列为管制了。以往各国数论家无忧无虑发表论文自由交换意见的日子也许是一去不返了，前程固然似锦，往事却已如烟，纯数学最后一块净土也终于被实用所“污染”了。然而这次因密码而推动大家对数论的研究，将在数学史上写下有趣的一页。

这种密码所用的数论并不深，我们可以全部介绍出来，当然在实际用的时候，数学会大得多，但在大小型电子计算机如此普遍的今日，是不会成问题的。

在我们介绍两种主要的数论密码之前，我们先将介绍一点数论。

§ 2 因子、质数、同余数与费马尤拉定理

若 m 、 n 为两整数，且 $m > 0$ ，则以 m 除 n 可得两个整数 d 与 r ，使得

$$n = dm + r \quad (0 \leq r < m)$$

其中 d 称为不完全商， r 称为余数。若 $r = 0$ ，则我们说 n 可被 m 所整除， m 为 n 之因子， n 为 m 之倍数，若一数除1与本身之外无其他因子，则此数称为质数，例如2，3，5，7，11都是质数，4，6，9，12却不是质数。我们定义 n 与 s 对 m 有同余数：

$$n \equiv s \pmod{m}$$

是指如果 n 与 s 被 m 除时有相同的余数。例如

$$12 \equiv 2 \pmod{10}$$

$$8 \equiv 5 \pmod{3}$$

有一个关于同余式的简单定理，我们把它们列出来，读者很容易证出来。

定理2.1 若 $p \equiv q \pmod{m}$

$$a \equiv b \pmod{m}$$

$$\text{则 } ap \equiv bq \pmod{m}$$

若两正整数 p 、 q 的最大公因子(约数)是1，则我们称 p 、 q 互质，以

$$(p, q) = 1$$

表示之。现在我们要证一个有关两个互质数的一个基本定理。

定理2.2 若两正整数 p 、 q 互质，则可以找到二整数(不一定正) a 、 b ，使得

$$ap + bq = 1$$

证明：令 A 为含所有 $x = ap + bq > 0$ ， a 、 b 为整数之集合，此集合显然不是空集合，因可取 $a = b = 1$ ， $p + q > 0$ 。令 d 为此集合中之最小者，若 $d = 1$ ，则本定理得证，若 $d > 1$ ，令 $ap + bq = d > 1$ ，则任取此集中之另一数 $a'p + b'q$ ，则我们若以 d 除 $a'p + b'q$ 有

$$a'p + b'q = ad + r \quad 0 \leq r < d$$

代入 $d = ap + bq$

则得

$$(a' - aa)p + (b' - ab)q = r$$

此 r 必为0，否则 r 为 A 集中一小于 d 之数，与假设 d 为最小数相矛盾，因 $r = 0$ 故 d 为 A 集中任何一数之因子。因

$$p \in A \quad (a = 1, b = 0)$$

$$q \in A \quad (a = 0, b = 1)$$

故 d 为 p 、 q 之公因子，但 p 、 q 之最大公因子为1，故 $d = 1$ 。定理证毕。

这是一个极有用的定理，读者也许要问，我们如何找到 a 与 b 使 $ap + bq = 1$ 呢？一般可用辗转相除法。

例：找整数 a 、 b ，使得 $5a + 9b = 1$ 。因

$$9 = 5 + 4, \quad 5 = 4 + 1$$

$$\text{故} \quad 1 = 5 - 4 = 5 - (9 - 5)$$

$$= 2 \times 5 - 9 \times 1$$

$$\text{故} \quad a = 2, \quad b = -1$$

推论2.1 若 w 与 m 为二互质的正整数且 $m > w$,则可找到一正整数 θ 使得

$$w\theta \equiv 1 \pmod{m}$$

证明: 由定理知, 有 a 、 b 二整数使得 $aw + bm = 1$, 因 bm 为 m 之倍数, 故

$$aw \equiv 1 \pmod{m}$$

$$\text{令 } a = \phi m + \theta \quad 0 \leq \theta < m$$

$$\text{则得 } \theta w \equiv 1 \pmod{m} \text{ 且 } \theta \geq 0$$

因 θ 不可能为0, 故本推论得证。

最后我们要用到一个不容易证明的“费马、尤拉(Fermat-Euler)定理”。但因为我们只用到这个定理比较容易证明的特殊形式, 我们就只证明简单的部分。

定理2.3 若 m 为质数, w 为任一与 m 互质的整数, 则

$$w^{m-1} \equiv 1 \pmod{m}$$

证明: 先把 w 写成 w 个1的和, 则由多项式定理知

$$(1 + 1 + 1 + \cdots + 1)^m$$

的展开式中除 w 个1之外, 都含有 m 的因子 (m 为质数,

$$m! = C_m^r = \frac{m!}{r! (m-r)!} \text{ 中之 } m \text{ 不可能消去), 故}$$

$$w^m \equiv w \pmod{m}$$

两边乘以推论1中的 θ , 即

$$\theta w \equiv 1 \pmod{m}$$

$$\text{得 } w^{m-1} \equiv 1 \pmod{m}$$

本定理证毕。

推论2.2 若 m 为两质数 p 、 q 之积, w 为任一与 m (即同时与 p 与 q) 互质之整数, 则

$$w^{(p-1)(q-1)} \equiv 1 \pmod{m}$$

证明: 先用定理2.3之证明法得

$$w^{p-1} \equiv 1 \pmod{p}$$

$$w^{q-1} \equiv 1 \pmod{q}$$

由定理2.1可得

$$(w^{p-1})^{q-1} \equiv 1 \pmod{p}$$

同理可得

$$(w^{q-1})^{p-1} \equiv 1 \pmod{q}$$

由上面两式及 p 、 q 互质, 不难得到

$$w^{(p-1)(q-1)} \equiv 1 \pmod{pq}$$

$$\text{即 } w^{(p-1)(q-1)} \equiv 1 \pmod{m}$$

本推论证毕.

我们还要利用到另一个简单的定理, 我们也在这一节里把它证完.

定理2.4 令 $(a_1, a_2, \dots, a_n)$ 为一含 n 个正整数的数列, 并满足

$$\begin{cases} a_1 \geq 1 \\ a_2 > a_1 \\ a_3 > a_1 + a_2 \\ \vdots \\ a_i > a_1 + a_2 + \dots + a_{i-1} \\ \vdots \\ a_n > a_1 + a_2 + \dots + a_{n-1} \end{cases}$$

又令 $(x_1, x_2, \dots, x_n)$ 为一由0与1组成的数列, 则所有的 x_i 不是0就是1, 现设 $a_1, a_2, \dots, a_n$ 为已知, $x_1, x_2, \dots, x_n$ 为未知, c 为一正整数, 则方程式

$$c = \sum_{i=1}^n a_i x_i \quad (2.2)$$

只有一解或无解。

证明：设此方程式有二解 $(x_1, x_2, \dots, x_n)$ 及 $(x'_1, x'_2, \dots, x'_n)$ ，则消去 c 之后可得

$$\sum_{i=1}^n (x_i - x'_i) a_i = 0$$

$$\text{因 } |x_i - x'_i| \leq 1$$

$$\text{故 } \left| \sum_{i=1}^{n-1} (x_i - x'_i) a_i \right| \leq \sum_{i=1}^{n-1} a_i < a_n$$

$$\text{由此推得 } x_n - x'_n = 0$$

$$\text{且 } \sum_{i=1}^{n-1} (x_i - x'_i) a_i = 0$$

同理可得

$$x_{n-1} = x'_{n-1}, \dots, x_1 = x'_1$$

此两解原为一。本定理得证。

而要解 (2.2) 是非常容易的事，因为若

$$c > a_1 + a_2 + \dots + a_{n-1}$$

则 x_n 必为1，否则必为0，同理若

$$c - x_n a_n > a_1 + a_2 + \dots + a_{n-2}$$

则 x_{n-1} 必为1，否则必为0，以此类推，一下子就解出来了。

这几个定理就足够解新的密码法了。

§3 狄飞、赫尔曼、麦克儿 (Diffie-Hellman、Merkle) 法

此法是由上列三位电机工程师兼数学家（科学本是一家，触类旁通，称他们是什么家并不正确，也不重要），在1976 及1978年所发表的方法。在此方法中，所有的信号必须写成二进位数的形式，假如前面的“我”字代码3314，若以二进位表示则为 $2^{11} + 2^{10} + 0 \cdot 2^9 + 0 \cdot 2^8 + 2^7 + 2^6 + 2^5 + 2^4 + 0 \cdot 2^3 + 0 \cdot 2^2 + 2^1 + 0 \cdot 2^0$ ，则

110011110010

我们将任何一位信号分成许多段，每段含 n 个0与1的数，即 $(x_1, x_2, \dots, x_n)$ ，以下是本密码法的收发程序（参看图5）。

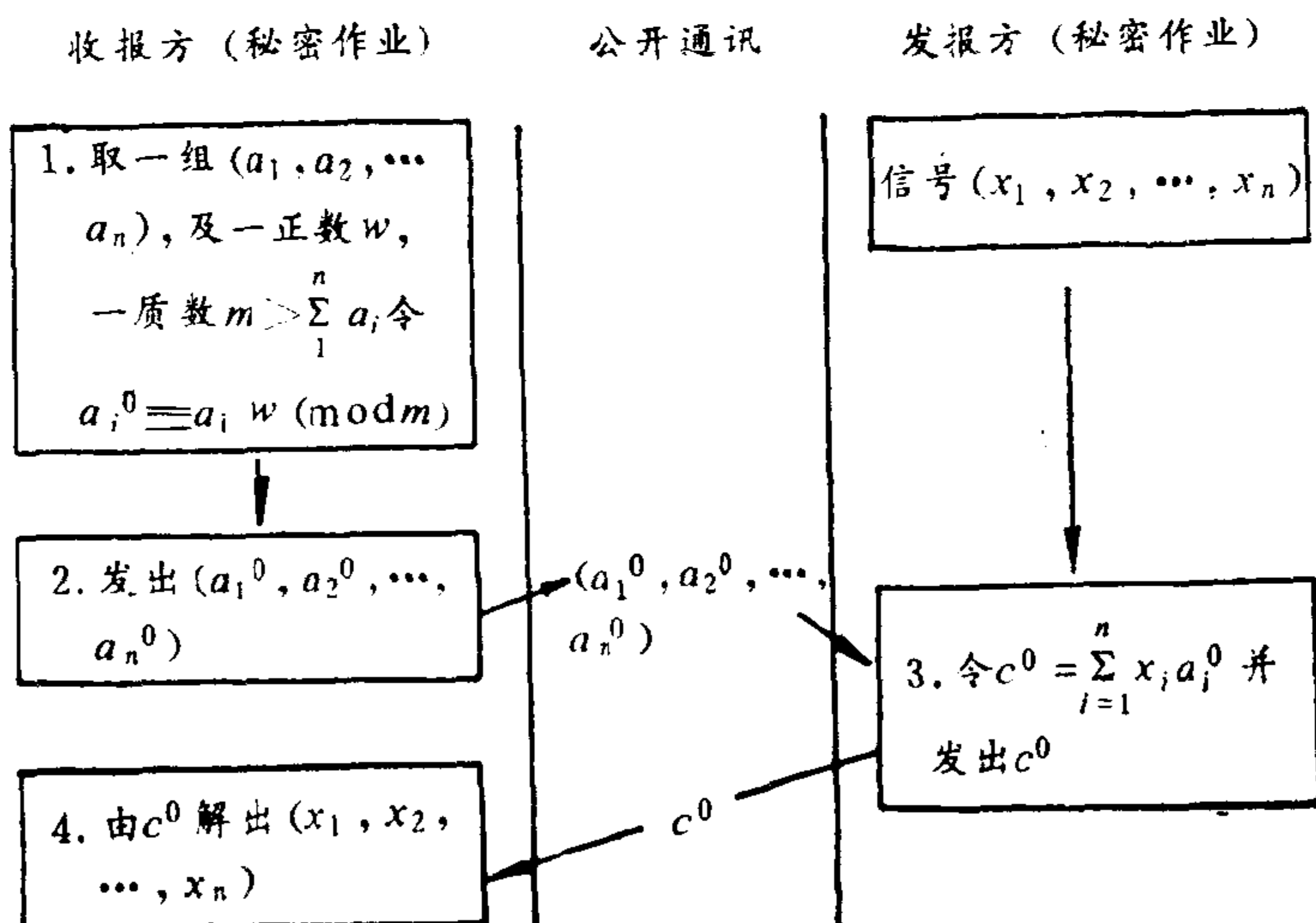


图5 在此种操作中，敌方仅知 $(a_1^0, \dots, a_n^0)$ 及 c^0 ，而发方除了自己的信号外并不比敌方多知道什么密码的机密；只有发方完全掌握了解码的方法。

1. 由收报者秘密的取一组满足定理2.4的正整数列 $(a_1, a_2, \dots, a_n)$ 及任一大于 $\sum_{i=1}^n a_i$ 的质数 m , 及另一数 w , 令

$$a_i^0 \equiv a_i w \pmod{m}$$

2. (公开)发出 $(a_1^0, a_2^0, \dots, a_n^0)$ 给发报者 (敌方可以知道).

3. 发报者用 $a_1^0, a_2^0, \dots, a_n^0$ 与要发的0, 1信号 $x_1, x_2, \dots, x_n$ 求出

$$c^0 = \sum_{i=1}^n x_i a_i^0$$

并将 c^0 之值告诉发报者 (敌方可以知道 c^0).

4. 收报者收到 c^0 之后, 可以解出原信号 $x_1, x_2, \dots, x_n$ (而敌方却不容易用已知的 $a_1^0, a_2^0, \dots, a_n^0$ 及 c^0 解出信号, 原因马上会谈).

我们先看收方如何解出 x_i , 由推论1可知存在 $-\theta > 0$ 且 $w\theta \equiv 1 \pmod{m}$, 若将收到之信号方程式两边乘以 θ 可使

$$c^0 = x_1 a_1^0 + x_2 a_2^0 + \dots + x_n a_n^0 \quad (3.1)$$

变成

$$c^0 \theta = x_1 a_1^0 \theta + x_2 a_2^0 \theta + \dots + x_n a_n^0 \theta \quad (3.2)$$

但 $a_i^0 \theta \equiv a_i w \theta \equiv a_i \pmod{m}$

故令

$$c^0 \theta \equiv c \pmod{m}$$

(3.2)即成为

$$c \equiv x_1 a_1 + x_2 a_2 + \dots + x_n a_n \pmod{m}$$

因 $m > \sum_{i=1}^n a_i$, 上式与 $c = x_1 a_1 + x_2 a_2 + \dots + x_n a_n$ 相同, 根据定

理2.4之说明，一下子就可以解出 $x_1, x_2, \dots, x_n$ ，可是敌方在整个的过程中知道(3.1)之关系，因 a_i^0 并不见得是一个有 a_i 那样规则的数列，到目前为止只有硬试一条途径，即

$$(x_1, \dots, x_n) = (1, 0, 0, \dots, 0), (0, 1, 0, \dots, 0) \dots$$

等一个一个地试，对小的 n 并不难，但对大的 n 而言，比如说 $n=1000$ （这并不是一个很长的信号），则平均要试 $2^{1000-1} = 10^{301}$ 次才可以解开，这是一个天文数字，若电脑在一秒钟内可以做一百万次检验（目前尚办不到），则一年只可以做 3.15×10^{13} 次检验，那么要解开(3.1)需要 10^{288} 年，前面谈过地球的年龄不过 5×10^9 年而已，我们且举一个简单的例子。

$$\begin{aligned} \text{例：令 } n=5, \quad (a_1, a_2, a_3, a_4, a_5) \\ = (1, 3, 6, 12, 25) \end{aligned}$$

$$w=13, \quad m=51$$

则收方发出的密码法 $(a_1^0, a_2^0, a_3^0, a_4^0, a_5^0)$ 分别为

$$a_1^0 \equiv 13 \times 1 \equiv 13 \pmod{51}$$

$$a_2^0 \equiv 13 \times 3 \equiv 39 \pmod{51}$$

$$a_3^0 \equiv 13 \times 6 \equiv 27 \pmod{51}$$

$$a_4^0 \equiv 13 \times 12 \equiv 3 \pmod{51}$$

$$a_5^0 \equiv 13 \times 25 \equiv 19 \pmod{51}$$

这即是发报者收到的作密码的 a^0 ，假定发方所要发的情报是10101，则因

$$c^0 = 13 + 27 + 19 = 59$$

故收方所收到的密码是59，要解开此码，收方先得找到 θ 使 $\theta w \equiv 1 \pmod{51}$ 。用辗转相除法很快得到 $\theta=4$ ，故收方所要解的是

$$c \equiv 59 \times 4 \equiv 32 \pmod{51}$$

$$\text{即} \quad 32 = x_1 + 3x_2 + 6x_3 + 12x_4 + 25x_5$$

$$= 1 + 6 + 25$$

即原信号为10101, 至于敌方所要解的是

$$59 = 13x_1 + 39x_2 + 27x_3 + 3x_4 + 19x_5$$

这自然不算太难, 但我们可以看见 a_i^0 失去了任何规则, 当 n 很大时就不容易解开了。

§4 瑞未斯特、希米尔、爱得曼 (Rivest, Schamir, Adleman) 法

这个方法是上列三位科学家在1978年所发表的, 其步骤如下 (如图6):

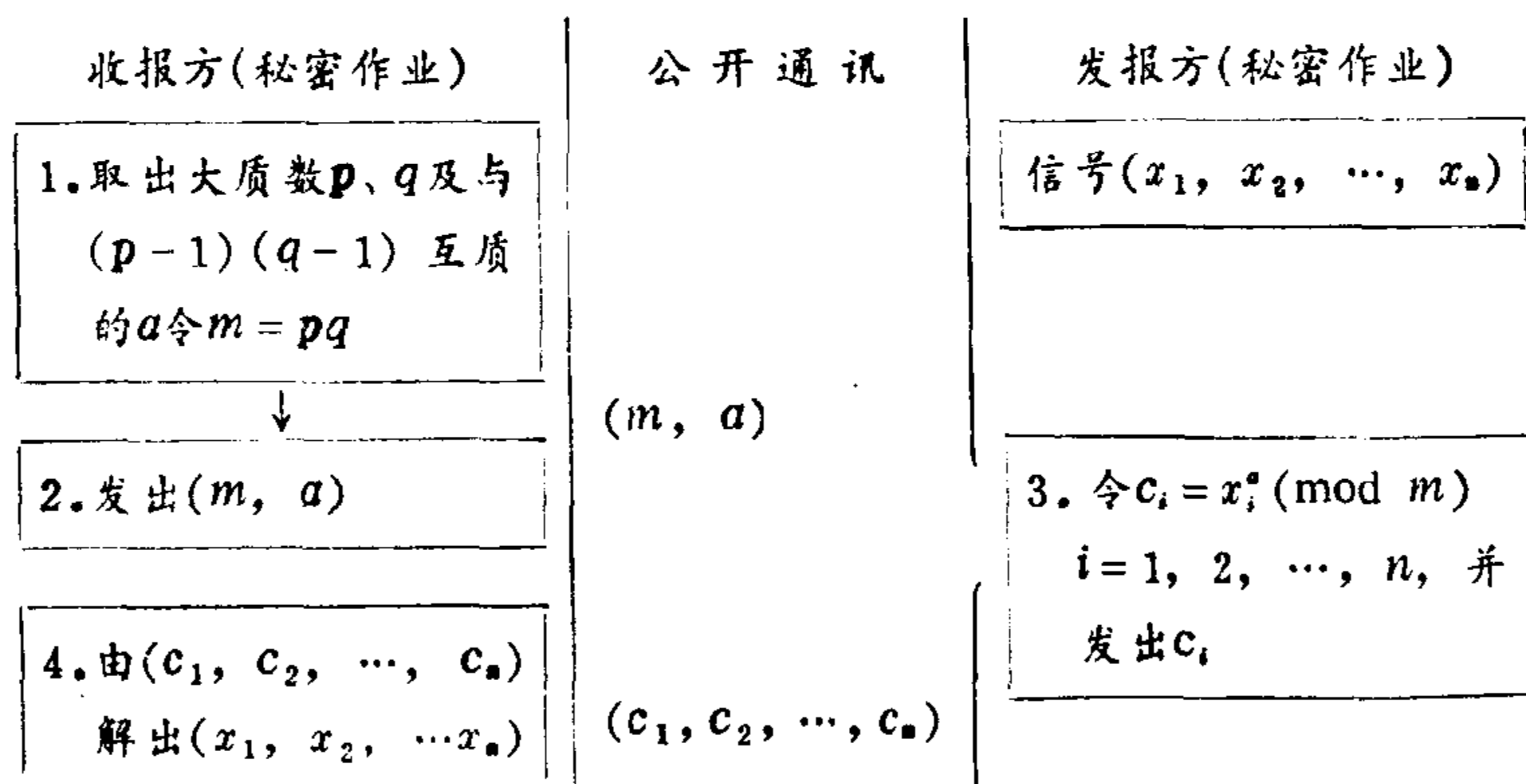


图6 (RSA制码法)在此作业程序中, $m, a, c_1, \dots, c_n$ 是公开的

1. 收报者取两个相异的大质数 p, q 及另一与 $(p-1)(q-1)$ 互质的数 a , 且 $a < w$, 令

$$w = (p-1)(q-1), \quad m = pq$$

及 p, q 之较小者的位数(十进位)为 K 。

2. (公开)告诉发报者 k 、 m 与 a 。

3. 发报者将他的信号分成许多段，每段含 $k-1$ 位数（十进位），若 $k=3$ （即 p 、 q 均为不小于二位的数），则信号

331414320001

应分成

33, 14, 14, 32, 00, 01

(一个一个的考虑发出)，设发报者的信号之一为 x ($k-1$ 位数，即上例中的33，或14，或32， $\dots$)，则他将它作成

$$c \equiv x^a \pmod{m}$$

发出。

4. 收报者收到 c 之后，即可把原来的 x 求出来，因 a 与 w 互质，由定理2.2及推论1知，我们可找到二整数 d 、 e ， $d > 0$ 使得

$$ad + we = 1$$

令

$$y \equiv c^d \pmod{m}$$

则此 y 即发报者之 x 。我们先证明 $y = x$ 。

$$y \equiv c^d \equiv (x^a)^d \equiv x^{ad}$$

$$\equiv x^{1-we}$$

$$\equiv x(x^{-we}) \pmod{m} \quad (4.1)$$

但因 w 、 a 、 d 均为大于1之整数，故 e 必为一负数，即 $-e$ 为一正数，又因 x 小于质数 p 、 q ，故 x 同 m 互质，由推论2得

$$x^w \equiv x^{(p-1)(q-1)} \equiv 1 \pmod{m}$$

故(4.1)或成为

$$y \equiv x(x^w)^{-e} \equiv x \cdot 1 \equiv x \pmod{m}$$

但因 y 与 x 均取小于 m 之数，故 $y = x$ ，故本程序之正确性得证。

这种密码的关键在于 p 、 q 为两大质数时，分解 m 成为 p 、 q 为一件极费时的工作，若分解不开 m ，则找不到 w 与 d ，因此就

无法从 c 解得 x 。在不久以前，要分解一个数的因子仍停留在近乎硬试的阶段，即要从2, 3, 5, 7..., 一直试到 $\sqrt{n}$ 附近才停止。若 n 是50位数而 p 、 q 均近25位数，则分解 m 要除约 10^{25} 次，若以电子计算机以每秒 10^6 次的高速运算，这仍是一个 10^{11} 年的工作，目前由于大家对这方面的重视，分解一个50位数的时间已可缩短至 10^{10} 次运算。下面的表中列出了目前(1980年)，分解一个大数大概所需的时间。

m 的位数	分解 m 的最少运算次数	最快(1980年)电脑所费的时间
50	1.4×10^{10}	3.9小时
70	9.0×10^{12}	104天
80	1.3×10^{18}	150天
100	2.3×10^{17}	74年
200	1.2×10^{23}	3.8×10^9 年

若取 p 、 q 各为40位数在目前已经十分安全了，即使是25位数，在商业上也十分安全，因为3.9小时最快电脑的计算费用也是一笔大的财富。

读者也许要问这种40位左右的大质数是否很多，而且容易找到否？答案是：这类大质数既多而且不难找（是指可用电脑很快制造出来），前面已谈到找一个质数不宜用硬除的办法，但因找它们所用的定理比我们已证明的几个要难，我们将在下节中才介绍出来，我们只谈一下质数的分布。依据质数定理在1与 n 之间的质数的有 $n/\ln n$ 个，因此小于 10^{40} （即40位）的质数约有

$$10^{40}/\ln 10^{40} = \frac{10^{40}}{92.1} \geq 10^{38} \text{ 个}$$

这又是一个天文数字，因为一个一千顿电子计算机中所言的原子数不过 10^{31} 左右，可见这种密码之难以捉摸了。现取两个用来做密码的三十位质数以飧读者：

$$p = 5862031427 \quad 1421210354 \quad 0772438083$$

$$q = 7976488510 \quad 8326808223 \quad 7297393713$$

要分解

$$m = pq$$

$$= 4675842632 \quad 8739231725 \quad 4879360844$$

$$\times 8514393251 \quad 3976539392 \quad 0565972179$$

若不懂数学与计算机，则是谈何容易。在本节结束之前，我们也举一个例子，当然我们不会用上列的大质数，我们且取 p, q 均为两位数，令

$$p = 47, \quad q = 59$$

则 $m = pq = 47 \times 59 = 2773$

$$w = (p - 1)(q - 1) = 2668$$

取 $a = 157$ ，由辗转相除可得

$$17a - 1w \equiv 1$$

故 $d = 17$

故收方发出的密码法是

$$m = 2773, \quad a = 157, \quad k = 2$$

此时发方必须一位一位的发出信号，设第一个要发的信号是3，则他要发出的是

$$c \equiv x^a \equiv 3^{157} \pmod{2773}$$

c 之求法主要在将157分成2进位数并用定理2.1，即

$$x^2 \equiv (x \pmod{m})^2 \pmod{m}$$

因 $157 = 2^7 + 2^4 + 2^3 + 2^2 + 1$ ，而

$$3 \equiv 3 \pmod{2773}$$

$$3^2 \equiv 9 \pmod{2773}$$

$$3^{2^2} \equiv 9^2 \equiv 81 \pmod{2773}$$

$$3^{2^3} \equiv 81^2 \equiv 1015 \pmod{2773}$$

$$3^{2^4} \equiv 1015^2 \equiv 1442 \pmod{2773}$$

$$3^{2^5} \equiv 1442^2 \equiv 2387 \pmod{2773}$$

$$3^{2^6} \equiv 2387^2 \equiv 2027 \pmod{2773}$$

$$3^{2^7} \equiv 2027^2 \equiv 1916 \pmod{2773}$$

故

$$\begin{aligned} 3^{167} &\equiv 1916 \times 1442 \times 1015 \times 81 \times 3 \pmod{2773} \\ &\equiv 964 \times 1015 \times 81 \times 3 \pmod{2773} \\ &\equiv 2364 \times 81 \times 3 \pmod{2773} \\ &\equiv 441 \pmod{2773} \end{aligned}$$

因此 $c = 441$ 即发方发出之信号，当收方收到 441 之后，用同样的运算法可得 c^d 为

$$\begin{aligned} 441^{17} &\equiv 441^{2^4+1} \pmod{2773} \\ &\equiv 371^{2^4} \times 441 \pmod{2773} \\ &\equiv 1764^{2^3} \times 441 \pmod{2773} \\ &\equiv 390^2 \times 441 \pmod{2773} \\ &\equiv 3 \pmod{2773} \end{aligned}$$

解码完成。由上面的运算可知若没有电子计算机，则解与做这种密码是何等的辛苦。

§ 5 如何寻找大质数

现在也许你有兴趣学一点难一些的数论定理了。

定理 2.5 设 a 、 m 互质，且

$$ab \equiv ac \pmod{m}$$

则
$$b \equiv c \pmod{m}$$

证明：因 $a(b-c) \equiv 0 \pmod{m}$ 而 a 不含 m 之因子，故 $b-c$ 必为 m 之倍数，即 $b-c \equiv 0 \pmod{m}$ ，本定理得证。

定理2.6 设 m 为一自然数， $\phi(m)$ 表示所有小于 m 且与 m 互质的自然数的数目，又以 $r_1, r_2, \dots, r_{\phi(m)}$ 表示这些自然数。如果 a 与 m 互质，则 $ar_1, ar_2, \dots, ar_{\phi(m)}$ 对 $\pmod{m}$ 而言是 $r_1, r_2, \dots, r_{\phi(m)}$ 的一个排列。

证明：令
$$x_i = ar_i \pmod{m}$$

且
$$0 \leq x_i < m$$

则
$$ar_i = qm + x_i$$

因 ar_i 与 m 互质，故 x_i 必与 m 互质，因此 x_i 是 $r_1, r_2, \dots, r_{\phi(m)}$ 中的一员，但由定理5.2知 $x_i \equiv x_j \pmod{m}$ 的充要条件为 $r_i = r_j$ ，故所有的 x_i 皆不相同，本定理得证。

定理2.7（费马、尤拉定理）设 w 与 m 为两互质的自然数， ϕ 为尤拉函数，则

$$w^{\phi(m)} \equiv 1 \pmod{m}$$

证明：由定理2.6知

$$\begin{aligned} & r_1 r_2 \cdots r_{\phi(m)} \\ & \equiv (wr_1)(wr_2) \cdots (wr_{\phi(m)}) \pmod{m} \end{aligned}$$

即
$$\begin{aligned} & r_1 r_2 \cdots r_{\phi(m)} \\ & \equiv w^{\phi(m)} r_1 r_2 \cdots r_{\phi(m)} \pmod{m} \end{aligned}$$

因 $r_1 r_2 \cdots r_{\phi(m)}$ 与 m 互质，故

$$w^{\phi(m)} \equiv 1 \pmod{m}$$

本定理得证。

（此定理为定理2.3的推广，因 m 为质数时 $\phi(m) = m-1$ ）

我们又可以推广定理2.2，若以 (a, b) 表示 a, b 的最大公

约数, 若 $d=(a, b)$, 则存在两整数 x, y 可使

$$ax + by = d$$

这个证法与定理2.2极相似, 求法也是用辗转相除法.

定理2.8 设 a, m 为两个互质的正整数且 $1 < a < m$, 若 $a^h \equiv 1 \pmod{m}$, 则 k 与 $\phi(m)$ 不互质.

证明: 设 A 为所有 $k \geq 1$ 且满足

$$a^k \equiv 1 \pmod{m}$$

的数的集合, 因 $\phi(m) \in A$, 故 A 为非空集合. 令 d 为 A 中之最小数, 因 $a^1 \equiv a \pmod{m}$, 故 $1 \in A$, 即 $d > 1$, 设 $a^h \equiv 1 \pmod{m}$, 令 $k = qd + r$, $0 \leq r < d$, 显然由 $a^d \equiv 1 \pmod{m}$ 可得

$$\begin{aligned} 1 &\equiv a^k \equiv a^{qd+r} \equiv (a^d)^q a^r \\ &\equiv a^r \pmod{m} \end{aligned}$$

但因 d 为此集合中的最小者, 故 $r = 0$, 即 k 必为 d 之倍数, 且因 $\phi(m) \in A$, 可知 $d \mid \phi(m)$, 从而 $(k, \phi(m)) \geq d > 1$, 故 k 与 $\phi(m)$ 不互质.

现在我们终于走到了最后一个求证某数是质数的定理.

定理2.9 (鲁克斯(Lucas)定理) 设 m 为大于1的整数, a 为任一与 m 互质之数. 若 $a^{m-1} \equiv 1 \pmod{m}$, 且对 $(m-1)$ 所有的因子 k 而言, $a^k \not\equiv 1 \pmod{m}$, 则 m 必为一质数.

证明: 设 m 不是质数, 则由 $\phi(m)$ 之定义可知 $\phi(m) < m-1$, 但由于 a, m 互质, 由定理2.7知 $a^{\phi(m)} \equiv 1 \pmod{m}$, 由已知条件及前定理知

$$(m-1, \phi(m)) = d > 1$$

故 d 为 $m-1$ 之一因子, 由(5.1)知存在二整数 x, y , 使得

$$(m-1)x + \phi(m)y = d$$

即

$$\begin{aligned} a^d &\equiv a^{(m-1)x + \phi(m)y} \\ &\equiv (a^{m-1})^x \cdot (a^{\phi(m)})^y \end{aligned}$$

$$\equiv 1(\text{mod } m)$$

故 $(m-1)$ 有一因子 d 使得 $a^d \equiv 1(\text{mod } m)$ ，与假设相矛盾，故 m 必为质数，本定理证毕。

这个定理若用来检定一已知数 m 是否是质数并不容易，因为 $m-1$ 的因子往往不易找到，前面已谈过要分解一个大数的因子是件极费时的事，但这个定理来找一些大质数却很容易，例如我们令

$$m = 2^n + 1$$

则 $m-1$ 的因子为 $2, 2^2, \dots, 2^n$ ，若我们能找到一个与 m 互质的数 a ，像 $3, 5, 7$ 之类的而且证得

$$a^{m-1} \equiv 1(\text{mod } m)$$

但 $a^{2^{n-1}} \not\equiv 1(\text{mod } m)$

则 m 必为质数。原因是若

$$a^{2^{n-1}} \equiv 1(\text{mod } m)$$

则对所有的 $k < n-1$ ， a^{2^k} 都不等于 $1(\text{mod } m)$ 。（假如 $a^{2^k} \equiv 1(\text{mod } m)$ ，则

$$a^{2^{n-1}} \equiv [a^{2^k}]^{2^{n-1-k}} \equiv 1(\text{mod } m)$$

对一个固定的 a 而言，这个整个检验过程不过一次验定 a 与 m 互质，二次 $(\text{mod } m)$ 的计算而已，比硬除快多了。

例：若令 $m = 2^n - 1$ ，则我们若能试出一与 m 互质的数 a 像 $3, 5, 7, \dots$ 而也证得

$$a^{m-1} \equiv 1(\text{mod } m)$$

且 $a^2 \not\equiv 1(\text{mod } m)$

及 $a^{2^{n-1}-1} \not\equiv 1(\text{mod } m)$

则 m 必为一质数（因 $m-1 = 2^n - 2 = 2(2^{n-1} - 1)$ ）。目前最大的

质数多半都是这样求得的，在1979年纳尔逊与斯罗温斯基 (H. Nelson & Slowinski) 用电子计算机证明 $2^{44497} - 1$ 是一个13395位的质数。若以100元一千字作稿费，把这个数写出来就是一千三百元，《数学传播》一页大约二千字，所以把这个字写出来要占去六页半的篇幅。

例：若 p 为一质数， s 、 t 为两正整数， $m = 2^s p^t + 1$ ，则我们若能试出一个与 m 互质的数 a 并证得

$$a^{m-1} \equiv 1 \pmod{m}$$

且 $a^{2^{s-t} p^t} \not\equiv 1 \pmod{m}$

$$a^{2^s p^{t-1}} \not\equiv 1 \pmod{m}$$

则 m 为一质数。

例：若 p 、 q 为两质数， $m = 2pq + 1$ ，则我们若能试出一个与 m 互质的数 a 并证得

$$a^{m-1} \equiv 1 \pmod{m}$$

$$a^{2p} \not\equiv 1 \pmod{m}$$

$$a^{pq} \not\equiv 1 \pmod{m}$$

及 $a^{2q} \not\equiv 1 \pmod{m}$

则 m 为一质数。

如此这般，我们可以由小而大滚雪球式的很容易找到许多大质数，由于走的路程千变万化，所可能找到的大质数也是一个天文数字。当然电子计算机是不可少的工具。在定理2.9中我们只要求任何一个与 m 互质的 a ，因此什么与 m 互质的 a 都可以，根据一般的经验，若 m 是质数，很快就可以找到一个小的 a 像3、5、7之类的满足定理2.9，若不成功，就可以放弃另找了，反正候选的 m 多得不得了，最近有一个极巧妙的结果由lehmer, Soloorary与Shassen同时发现。即若 m 不是质数，则至

少有一半以上的数从2, 3, ..., 到 $m-1$ 不能满足 $a^{m-1} \equiv 1 \pmod{m}$, 可惜我们没有办法在本文中证明此一结果。这个结果在另一个角度来看, 即一个不是质数的 m 只有很小的机会能通过多少的 a 。我们最后举一个数字的例子。

例: $257 = 2^8 + 1$ 是不是质数?

因 $m-1 = 2^8$, 故我们若能找到一个与 m 互质的 a , 且

$$a^{2^k} \equiv 1 \pmod{m}$$

$$a^{2^{k'}} \not\equiv 1 \pmod{m}$$

m 即为质数。 $m = 257$ 不是3的倍数, 我们先取 $a = 3$, 我们一步步的求 $3^{2^k} \pmod{m}$ 与 $3^{2^{k'}} \pmod{m}$ 。

k	2^k	$3^{2^k} \equiv ? \pmod{257}$
0	$2^0 = 1$	$3 \equiv 3$
1	$2^1 = 2$	$3^2 \equiv 9$
2	$2^2 = 4$	$3^4 \equiv 9^2 \equiv 81$
3	$2^3 = 8$	$3^8 \equiv 81^2 \equiv 136$
4	$2^4 = 16$	$136^2 \equiv 249$
5	$2^5 = 32$	$249^2 \equiv 64$
6	$2^6 = 64$	$64^2 \equiv 241$
7	$2^7 = 128$	$241^2 \equiv 256$
8	$2^8 = 256$	$256^2 \equiv 1$

故257是一个质数。

§ 6 结尾的话

前面说过, 在重赏或生死交关的情形下, 至今尚很少不能破获的密码, 在现今数论密码一日千里的进展之下, 我们刚才

谈到两种密码也许就要（已经？）落伍，但新的方法必然又会出现，并且在寻找新数论密码的长途上，数学家一定可以有时在路边捡到一些前人所未发现的珍珠。有一天人类也许可以坦诚相处到不要用密码通讯的地步，但这些拾来的珍珠将永远是数学史上的光辉。

补充资料

①据1985年3月份的科学新知 (Science News) 的报导于1985年2月间荷兰的阿姆斯特丹大学的一位数学家H. W. Lenstra Jr. 发表了仅有一页的一篇有关分解大数法则的研究摘要，惊动了数学界。据一些读到详文的专家评论此方法时认为，一般说来，是要比目前其它的方法来得快速有效率，而且原理也很简单，并称对此方面有素养的专家，只需几分钟就可把此方法加以描述出来，目前不少人正埋头写有关此法则的电脑运算程序，以便进行试验，也有不少人忙着发掘此法则的应用及用途。

此新法据称与具 $y^2 = x^3 + ax + b$ 形式的椭圆形曲线有关，式中 a 、 b 可随机地取定，但此新法不是对所有的大数都最有效，它对一个原由三个或三个以上质数相乘积的大数，或两个相差很大的质数相乘的大数最有效，因而此新法对我们上面介绍的“公开通讯”的密码仍不至于构成什么大威胁（因在该密码中，我们可取两个相差不大的大质数作我们的大数。当然理论上讲，这多少给了敌方一些线索了），但人们已开始对此密码的绝对安全可靠动摇了，希望在读者中能有人很快地将此忧虑造成事实，为中国人带来灿烂的光辉，但要有进展最好从数论及电脑程序设计（软硬件两方面）双方的配合着手，才能得到真

正有效用的新法则。

②在6月25日华府邮报(Washington post)科学栏的密码报导两位年青的数论专家，一位是贝尔实验室的蓝司特拉(A. Lenstra)及数据设备公司(Digital Equipment corp)的马拿塞(M. Manasse)在世界各地二百多个研究员的协助下(主要是利用他们的电脑来计算，约1千架电脑的空间时间来进行运算)花了几个月的时间，把以前一直认为无法去分解的一个具155位数分解出来了(所谓无法分解是指以当时的一架计算机的计算量来计有生之年是完成不了的)。

13,407,807,929,942,597,099,574,024,998,205,846,
127,479,365,820,592,393,377,723,561,443,721,764,030,
073,546,976,801,874,298,166,903,427,690,031,858,186,
486,050,853,753,882,811,946,569,946,433,649,006,084,
097

$= 2,424,833 \times 7,455,602,825,647,884,208,337,395,736,$
 $200,454,918,783,366,342,657 \times 741,640,062,627,530,$
 $801,524,787,141,901,937,474,059,940,781,097,519,$
 $023,905,821,316,144,415,759,504,705,008,092,818,$
 $711,693,940,737$

而目前一般商用、外交或军事上靠电话线传递数值化的分开密码也不过是用155位的数来制造的，而且一直认为至少在本世纪结束前以目前一架电脑计算量来衡量是可保险的。所以上面的一个突破的确曾引起了一些使用密码者的担忧。好在上面的155位数是属于一种特殊的数，即所谓的“费马数”(见形式 $2^{2^n} - 1$ 的数)。而且该两位数学家也承认他们的方法对别种形式的大数分解不见得有效。

③MIP一年及分解大数的计算估计。

首先我们称一部电脑以每秒进行一百万次加减运算，不断进行一年的计算量为一个MIP年。

现我们来介绍一下如何对大数的分解的时日有所估计。依据 R. Rivest 在他的一篇投书 (Science, Dec. 16, 1988, P. 1493)中指出，一般对一个 n 位数如照当时最好的运算法来分解，其计算量为 $\exp(\ln(n)\ln\ln(n))^{1/2}$ 。特别对一个100位数的分解需要25个MIP年。一般说来位数在100到300位间的大数，每增加一位，则其分解的计算量的增加20%。

也因有此估计，所以公开性的数字密码的保密性是相对制码的位数 n 与电脑计算量的改进而言的。

参考资料

1. Rivest, R. L; Schmir, A. ; Adleman, L. "A method for obtaining digital signature and public-key cryptasystem" Communication of ACM, 1978, pp. 120—126.
2. Simmons, G. "Cryptology: The mathematics of secure communication" The Mathematical Intelligencer, 1978, pp. 233—246.
3. Hellman. M. E. "The mathematics of public-key cryptography" Scientific American, August, 1978, pp. 146—157.
4. Pomerance, C. "The search for prime numbers" Scientific American, December, 1982. pp. 136—147.

第三章 数字密码的一些新研究

§1 前 言

在前文中我们介绍了Rivest-Schamir-Adleman在1978年发表了一种破天荒的数字密码。它和往常密码不同之处是，拍发及接受两方都不必怕拍出的数字密码中途被人截接侦破，只要收方保有一组解码的数字组，用来把拍出的数字转换成原来未拍出前的数字即可，其中所用的原理是整数论中的所谓“中国人剩余定理”，也是国人所称的“韩信点兵术”，及若仅知一个大数 N ，其为两个大质数 p_1 、 p_2 的乘积（即 $N = p_1 p_2$ ），但要把这两个质因子 p_1 、 p_2 找出是很困难的事实（一般如果 N 为一个百位数，照目前已知找质数的方法及最快速的计算机来计算，日夜不停地进行计算，也至少要花上九年的时间），所以上面提的密码算是目前最具保密性的了。但在该法中，拍码及译码（解码）时都用到了大数的高幂次计算，所以计算量很大，这是一个缺点（但也确保了保密性）。

在1979年两位华裔的工程师S. C. Lu及L. N. Lee, 介绍了一个较D-H两氏方法简便得多的一种密码[1], 它也是利用中国人剩余定理及找大质因子困难的事实, 但发码时只涉及几个乘法及对模数的加法, 译码时只涉及解一组二元一次方程式. 虽然Lu-Lee两氏在他的文中声称具有高度的保密性, 但在文[1]发表的同年就有人撰文在[3]及[4]中指出Lu-Lee两氏密码的破绽.

1985年两位印度工程师提出了一篇对Lu-Lee两氏方法的修正文章[2], 并且声称此法至少可以抵挡[3]及[4]两文中的破法. 我们在本文中先将介绍Lu-Lee两氏的密码法及其破绽之处, 然后介绍、讨论[2]文中所作的改进, 最后我们把最近在数学界有关质数的检验法的进展当作补充资料作报导, 希望读者能了解最新的研究, 在此方面作急起直追的努力.

§ 2 LU-Lee密码方法及原理

除非特别声明外, 本文用来代表未知数及已知数的文字都是指的整数. 我们将就拍码(或编码)法及解码(或译码)法作解说.

拍码法:

设 m_1 及 m_2 为两个要送出的数字码(或信息), 我们可不妨限定 $0 < m_1 < M_1$ 及 $0 < m_2 < M_2$, M_1 及 M_2 为有关数码 m_1 及 m_2 的上界. 以下我们要指出如何定此上界(这与解码有关). 我们可以公开一组与拍码有关的数字组 (c_1, c_2, r) .

现所要拍出的解码为 x ,

$$x \equiv (c_1 m_1 + c_2 m_2) \pmod{r} \quad (1)$$

解码法:

收方必须具有一组可以由 x 得出 m_1 及 m_2 的解码数字组
 $(a_{11}, a_{12}, a_{21}, a_{22}, p_1, p_2)$, 此组为保密的。

则当甲方拍出 x 时, 乙方作如下的计算:

$$x_1 \equiv x \pmod{p_1}, \quad x_2 \equiv x \pmod{p_2} \quad (2)$$

得了 x_1 及 x_2 后, m_1 及 m_2 可由下面公式得出:

$$m_1 = \frac{x_1 a_{22} - x_2 a_{12}}{a_{11} a_{22} - a_{12} a_{21}} \quad (3)$$

$$m_2 = \frac{x_2 a_{11} - x_1 a_{21}}{a_{11} a_{22} - a_{12} a_{21}} \quad (4)$$

任何一解码必须要满足对于拍出不同的 x , 一定得出不同的 m_1 及 m_2 , 否则就混淆不知原来的数码了, 如何有此种保证? 这与参数 c_1, c_2, r , 及 $a_{11}, a_{12}, a_{22}, a_{21}, M_1$ 及 M_2 的选择条件有关, 以下我们就讨论如何选取这些参数。

选参数的条件:

1. c_1 及 r 、 c_2 及 r 皆为互质, r 为两个大质数 p_1, p_2 之乘积, 即

$$(c_1, r) = (c_2, r) = 1, \quad r = p_1 p_2 \quad (5)$$

并且进一步要求

$$c_1 + c_2 \geq r$$

现要保密的解码参数组 $(a_{11}, a_{12}, a_{21}, a_{22}, p_1, p_2)$ 参数都是要保密的, 但它们之间要满足下列关系:

$$\begin{aligned} 2. \quad a_{11} &\equiv c_1 \pmod{p_1}, \quad a_{12} \equiv c_2 \pmod{p_1} \\ a_{21} &\equiv c_1 \pmod{p_2}, \quad a_{22} \equiv c_2 \pmod{p_2} \end{aligned} \quad (6)$$

并且要求

$$a_{11} a_{22} - a_{12} a_{21} \neq 0 \quad (7)$$

这是很明显的—个要求。

3. M_1 及 M_2 要求满足

$$M_1 \leq \left[\frac{1}{2} \min \left\{ \frac{q}{a_{11}}, \frac{q}{a_{21}} \right\} \right] \quad (8)$$

$$M_2 \leq \left[\frac{1}{2} \min \left\{ \frac{q}{a_{12}}, \frac{q}{a_{22}} \right\} \right] \quad (9)$$

其中 $q = \min\{p_1, p_2\}$, $[y]$ 表 y 的整数部分。

现在我们看解码成立的过程:

对式(1)两边取模数 p_1 , 由同余的原理可得,

$$x_1 \equiv x \pmod{p_1} \equiv (c_1 m_1 + c_2 m_2) \pmod{p_1} \quad (10)$$

或 $x_1 \equiv x \pmod{p_1}$

$$\equiv [c_1 \pmod{p_1} m_1 + c_2 \pmod{p_1} m_2] \pmod{p_1}$$

$$\equiv (a_{11} m_1 + a_{12} m_2) \pmod{p_1} \quad (11)$$

由条件(8)及(9)可知 $a_{11} m_1 + a_{12} m_2 \leq p_1$, 因而

$$x_1 = a_{11} m_1 + a_{12} m_2 \quad (12)$$

同理可得

$$x_2 = a_{21} m_1 + a_{22} m_2 \quad (13)$$

解方程式(12)及(13)可得 m_1 及 m_2 , 如(3)式及(4)式, 此解的存在及唯一性是因为 $a_{11} a_{22} - a_{12} a_{21} \neq 0$, 又注意 m_1 及 m_2 皆为整数。

读者不难发现这个方法的计算量很少。

参数的选取:

大参数 p_1 及 p_2 的选取可依据 Solovay 及 Strassen[6] 所提出的有效的或然率选法。一旦 p_1 及 p_2 选了, 就可选 a_{11} , a_{12} , a_{21} 及 a_{22} 满足 $a_{11} a_{22} - a_{12} a_{21} \neq 0$, 然后由于 p_1 及 p_2 互质, 由 Euclid 的辗转相除法可得,

$$b_1 p_1 + b_2 p_2 = 1 \quad (14)$$

对上式两边同乘以 $(a_{21} - a_{11})$ ，经并项后可得

$$c_1 \equiv [(a_{21} - a_{11})b_1p_1 + a_{11}](\text{mod } r) \quad (15)$$

或

$$c_1 \equiv [(a_{11} - a_{21})b_2p_2 + a_{21}](\text{mod } r) \quad (16)$$

上式中 $r = p_1p_2$ 。同理可得

$$c_2 \equiv [(a_{22} - a_{12})b_1p_1 + a_{12}](\text{mod } r) \quad (17)$$

或

$$c_2 \equiv [(a_{12} - a_{22})b_2p_2 + a_{22}](\text{mod } r) \quad (18)$$

下面我们举一个实际计算的例子([1]文中的)。由此例子我们可得到一点端倪，为何此法易被侦破了（当然参数选取的种种限制也多少提供了一些线索）。

例子：取 $p_1 = 97$ ， $p_2 = 103$ ，及 $a_{11} = 3$ ， $a_{12} = 2$ ， $a_{21} = 5$ ， $a_{22} = 4$ 作为解码的参数组（这是要保密的），现 $r = p_1p_2 = 9991$ 及 $1 = 17 \times 97 - 16 \times 103$ 可得 $b_1 = 17$ ， $b_2 = -16$ ，于是发码的参数组 (c_1, c_2, r) （这是要公开的）中的 c_1 及 c_2 可得如下：

$$c_1 = 2 \times 17 \times 97 + 3 = 3301$$

$$c_2 = 2 \times 17 \times 97 + 2 = 3300$$

现我们能拍送的数字码 m_1 及 m_2 能有多大？

$$M_1 \leq \left[\frac{1}{2} \min\left(\frac{97}{3}, \frac{97}{5}\right) \right] = 9$$

$$M_2 \leq \left[\frac{1}{2} \min\left(\frac{97}{2}, \frac{97}{4}\right) \right] = 12$$

换句话 m_1 及 m_2 的选择不能分别大于9及12，现我们比方取 $m_1 = 7$ ， $m_2 = 5$ （或用二进位制 $m_1 = 0111$ 及 $m_2 = 0101$ ）。

拍发的明码为

$$\begin{aligned} x &\equiv [7 \times 3300 + 5 \times 3300](\text{mod } 9991) \\ &\equiv 9634 \end{aligned}$$

因而

$$x_1 = 9634(\bmod 97) = 31$$

$$x_2 = 9634(\bmod 103) = 55$$

解

$$3m_1 + 2m_2 = 31$$

$$5m_1 + 4m_2 = 55$$

得到 $m_1 = 7$ 及 $m_2 = 5$ ，即原来的数字信息。

在文[2]中Lu-Lee两氏也曾讨论了一些可能的破绽，不过他们总结只要 p_1 及 p_2 取得够大，及一般 a_{ij} 也取得相当大就可避免被侦破的危险，但 a_{ij} 一大， M_1 及 M_2 就要减小了，这个矛盾是在[2]文中所要对付的，下面我们就介绍此一改进的方法。

首先在[2]文中指出的是在[3]及[4]的两篇文中利用拍送出不同的明码，经过译码得到不同的原码的事实，Lu-Lee两氏的密码可以不必知道 p_1 及 p_2 ，照样可把 m_1 及 m_2 求出。另外在文[4]中指出由于当 $c_i (i=1, 2)$ 取 p_1 或 p_2 为模所得的剩余数 a_{ij} 值都很少（这是因为要求 $a_{i1}m_1 + a_{i2}m_2 < p_i, i=1, 2$ ），所以依此事实可以把 p_1 及 p_2 求出，之后可求得 a_{ij} 。

§3 Lu-Lee法的改进

参数选取：

如同Lu-Lee法一样，译码的秘密解码组为一组数 $(p_1, p_2, a_{ij}, i=1, 2, j=1, 2)$ 及公开的发码，参数组为 (c_1, c_2, r) ，但要求 a_{ij} 满足

$$(1) a_{12} > a_{22}.$$

$$(2) a_{21} > a_{11}.$$

(3) a_{ij} , $i=1, 2, j=1, 2$, 每个值不小于 2^{200} (或二进制制200位的数字)。

(4) 对于 M_1 及 M_2 我们要求 $M_1 \leq 2^{50}$, $M_2 \leq 2^{50}$ 。

假定我们固定取 p_1 及 p_2 皆为二进制下具有252位数的值(因而 r 为一在二进制表示具有504位数的值), 就可以使得(5)满足了。

发码:

(1) 首先我们要求发的数码 m 不大于 2^{199} (即在二进制下至多为一个199位数)。

(2) 任选一组整数 (m_1, m_2) , $m_1 \leq M_1$, $m_2 \leq M_2$, 拍下列明的数字码

$$m_e \equiv (m + c_1 m_1 + c_2 m_2) \pmod{r}$$

换句话说, 把Lu-Lee中的明码加上了一个因子 $m \pmod{r}$, 注意的是这时我们是一次送一数码, 解一数码。

我们首先看这个密码法会不会产生混淆, 即不同的原码 m 及 m' 其相应的 m_e 及 m'_e 是否可能会相同?

所以我们假设 $m \neq m'$ 看 $m_e = m'_e$ 可不可能?

今

$$\begin{aligned} m_e &\equiv (c_1 m_1 + c_2 m_2 + m) \pmod{r} \\ &\equiv \{(c_1 m_1 + c_2 m_2) \pmod{r} + m \pmod{r}\} \pmod{r} \\ &\equiv (x_e + m) \pmod{r} \end{aligned}$$

同理可得

$$m'_e \equiv (x'_e + m) \pmod{r}$$

及

$$\begin{aligned} m_e \pmod{p_1} &\equiv \{x_e \pmod{p_1} + m \pmod{p_2}\} \pmod{p_1} \\ &\equiv (x_1 + m) \pmod{p_1} \end{aligned} \quad (19)$$

$$m_e(\bmod p_2) \equiv (x_2 + m)(\bmod p_2) \quad (20)$$

$$m'_e(\bmod p_i) \equiv (x'_i + m')(\bmod p_i), i=1, 2 \quad (21)$$

于是由 $m_e = m'_e$ 两边取 p_i 为模数的余式相等下, 利用上面三式可得:

$$(x_i + m)(\bmod p_i) \equiv (x'_i + m')(\bmod p_i); i=1, 2 \quad (22)$$

即

$$(x_i - x'_i)(\bmod p_i) \equiv (m - m')(\bmod p_i); i=1, 2$$

因 m, m', x_i, x'_i 皆小于 $p_i, i=1, 2$, 故

$$x_i - x'_i = (m - m'); i=1, 2, \text{ 因而}$$

$$x_1 - x'_1 = 2n' - m = x_2 - x'_2 \quad (23)$$

注意 $|(m - m')| \leq a_{i1}, i=1, 2$ 及 $j=1, 2$, 依定义 x_i 及式(6)

$$\begin{cases} x_1 = a_{11}m_1 + a_{12}m_2 \\ x_2 = a_{21}m_1 + a_{22}m_2 \end{cases}$$

及

$$\begin{cases} x'_1 = a_{11}m'_1 + a_{12}m'_2 \\ x'_2 = a_{21}m'_1 + a_{22}m'_2 \end{cases}$$

由上面两组方程组及(23)可得

$$\begin{aligned} & a_{11}(m_1 - m'_1) + a_{12}(m_2 - m'_2) \\ & = a_{21}(m_1 - m'_1) + a_{22}(m_2 - m'_2) \end{aligned}$$

于是

$$(a_{11} - a_{21})(m_1 - m'_1) = (a_{22} - a_{12})(m_2 - m'_2) \quad (24)$$

由于上式两方必须为同号, 故 $m_1 - m'_1$ 或 $m_2 - m'_2$ 同为正或同为负。若同为正, 则依据式(23)、由式(24)可得:

$$\begin{aligned} m - m' = x_1 - x'_1 &= a_{11}(m_1 - m'_1) + a_{12}(m_2 - m'_2) \\ &> a_{11} + a_{12} \end{aligned}$$

此与 m 与 m' 之大小规定不符, 同样在 $m_1 - m'_1$ 及 $m_2 - m'_2$ 同为负时亦可得同样的矛盾, 所以若 $m \neq m'$, 则 $m_e \neq m'_e$ 。

§ 4 改进码的解码步骤

$$\text{计算 } m_{ei} \equiv m_e \pmod{p_i}; \quad i = 1, 2 \quad (25)$$

解下列联立方程组（未知数为 t_1 及 t_2 ，其解可能为有理数，不一定为整数）。

$$\begin{cases} a_{11}t_1 + a_{12}t_2 = m_{,1} \\ a_{21}t_1 + a_{22}t_2 = m_{,2} \end{cases} \quad (26)$$

$$\text{取 } k_1 = [t_1] \text{ 及 } k_2 = [t_2] \quad (27)$$

计算 $a_{i1}k_1 + a_{i2}k_2 = m'_{ci}; i=1, 2$ (28)

$$\text{取} \quad m = m_{\varepsilon i} - m'_{\varepsilon i} \quad (29)$$

我们现看上面每一步骤是否合理？主要是找出步骤(26)中 $m_{e,1}$ 与 $m_{e,2}$ 及检验步骤(29)。

任何一有理数 t 都可成式 $t = [t] + (t - [t]) = \text{整数部分} + \text{小数部分}$ 。所以不妨设

$$t_i = [t_i] + r_i / \Delta$$

$$\Delta = \begin{vmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{vmatrix} = a_{11}a_{22} - a_{21}a_{12}, \quad r_i \text{ 为一适当的有理数, 使}$$

得 r_i/Δ 为 t_i 之小数部分.

我们把步骤(26)中的方程组改写为下面的形式:

$$\begin{cases} a_{11}t_1 + a_{12}t_2 = x_1 + m \\ \quad \quad \quad \text{(即以 } x_1 + m \text{ 表 } m_{e1}, x_1 \text{ 为一参数)} \\ a_{21}t_1 + a_{22}t_2 = x_2 + m \\ \quad \quad \quad \text{(即以 } x_2 + m \text{ 表 } m_{e2}, x_2 \text{ 为一参数)} \end{cases}$$

以下我们将证明 x_i 为步骤(28)中的 m'_i ; $i=1, 2$. 这样步骤(29)亦成立, 原数码也就得到了. 解上面方程组得

$$\begin{cases} t_1 = (a_{22}x_1 - a_{12}x_2)/\Delta + m(a_{22} - a_{12})/\Delta \\ t_2 = (a_{21}x_1 - a_{11}x_2)/(-\Delta) + m(a_{21} - a_{11})/(-\Delta) \end{cases}$$

由于我们将希望看到的是 x_i 为 $m'_{e,i}(=a_{i1}k_1 + a_{i2}k_2)$ ，所以我们不妨设

$$\begin{cases} x_1 = a_{11}m_1 + a_{12}m_2 \\ x_2 = a_{21}m_1 + a_{22}m_2 \end{cases}$$

再证明 m_1 与 m_2 事实上分别等于 k_1 及 k_2 就得了。由上面方程组可得：

$$\begin{cases} m_1\Delta = a_{22}x_1 - a_{12}x_2 \\ m_2\Delta = -a_{21}x_1 + a_{11}x_2 \end{cases}$$

将此组结果代入 t_1 及 t_2 之值中得

$$\begin{cases} t_1 = m_1 + m(a_{22} - a_{12})/\Delta \\ t_2 = m_2 + m(a_{21} - a_{11})/(-\Delta) \end{cases}$$

由于参数选取的条件，可验证得 $m(a_{22} - a_{12})/\Delta$ 及 $m(a_{21} - a_{11})/(-\Delta)$ 皆为小数。

$$[t_1] = m_1 \text{ 及 } [t_2] = m_2$$

因而在解码的步骤(28)中所计算的为方程组

$$\begin{cases} x_1 = a_{11}m_1 + a_{12}m_2 \\ x_2 = a_{21}m_1 + a_{22}m_2 \end{cases}$$

中的 x_1 及 x_2

有了
$$\begin{cases} x_1 = m'_{e,1} \text{ 及 } x_2 = m'_{e,2} \\ x_1 + m = m_{e,1} \text{ 及 } x_2 + m = m_{e,2} \end{cases}$$

因而

步骤(29)也就自然成立了。

讨论问题

在发码步骤(2)中， $m_e \equiv (m + c_1m_1 + c_2m_2) \pmod{2}$ ，但

$m + c_1 m_1 + c_2 m_2$ 也可表成 $m + c_1 - c_2 + c_1(m_1^{-1}) + c_2(m_2 + 1)$, 后者相当于一个人要求发的码为 $m + c_1 - c_2$, 选的整数组为 $(m_1 - 1, m_2 + 1)$. 对收方来说他收到相同的值 m_e , 为何此改进法仍可用? 读者不妨试由值 m, a_{ij}, p_1, p_2 等的大小, 来补充证明 m_e 不可能同时以上列二种形式来产生.

§ 5 补充资料

为了提起读者研究的兴趣, 我们在此顺便介绍最近在美国科学杂志 (Science, vol. 231 1985, p. 452—453) 指导的有关质数检验法的一些最新发展.

过去十年以来, 出现了许多有关质数的检验, 但都有它们独特的短处, 所以迄今为止仍无一个十全十美的方法. 在1975年左右任职哈佛大学及希伯来大学的 M. Rabin 分别与伯克莱加州大学的 R. Solovay 及瑞士苏黎士 (Zurich) 大学的 V. Strassen 两人小组, 各自发现了一种或然率的质数检验法, 此检验法容易且快速. 所谓快速是指用计算机来照此方法的程序运算, 所需的时间是多项式时间 (Polynomial time 较具体的讲, 所需的四则运算的次数为 $\leq cn^k$, n 为所要检定的正整数, k 及 c 为两常数. 详情可参看下章). 但由于它的或然性, 使得一般人存疑, 也因此不敢照用此法. 这个方法使得当一个人用它去试一个数, 如果检验得出的答案是非质数, 他一定可以相信 (也即结论一定是正确的), 但如果得出是质数的答案, 那就不一定百分之百的可信了. 因此理论上这个方法不能用来检验哪个数为质数.

几乎在同时, 南加大的 G. Miller 发现了一个非或然性的质数检验法. 她的方法可以很快 (多项式时间) 地把一个原非质

数的数，正确地判断出来。但若一个数经检验得的结论是质数的话，是不是一定正确，却要看数学上有名的推广的黎曼氏假说(Extended Riemann Hypothesis)是否成立而定。因此假如该假说被证明是错误的话，许多由Miller氏判定的质数，事实就是非质数了。

所以自然就有人试着把这些或然的性质，或上面方法中与假说相连的缺陷从方法中除去。果然在1980年三个美国数学家L. Adleman (南加大)，R. Rumley及C. Pomerance (伊里诺大学)得到了一个相当快而且具果断性的检验法。他们的方法可以正确地判断一个数为质数或非质数，但其缺陷是计算速度不够快(超过多项式时间)，仍无法改进到多项式的时间。

最新的一个方法是由麻省理工学院的两位数学家Goldwasser及 Kilian 的或然性检验法(这里的或然性与前面提的或然性两者意义上不同)，这个方法很快(只需多项式时间)就可以把一个非质数的数正确地判定出来，对于质数也可以作正确地判断，但其理论上就是可能对一小部分(比例上而言)的质数，这个方法判断的时间不很快(即超过多项式时间)，也就是因为有这样一个可能很小的机会，运算时间较长，所以仍把这个方法归之于或然性的检验法。但有趣的是迄今为止还未有人找出一个需要超过多项式时间来判断的质数例子！而且如果真的找出如此一个质数，将会轰动数学界，因为这与数论上一个有关质数分布密度的有名的猜测有关。

在K-G两氏的方法及其它一些类似的方法中，主要与研究所谓的椭圆曲线上的整数点及其阶(rank)的高深研究有关，而这方面的研究，原来是与解决讨论下列的一个问题有关：

设 $p(x, y)$ 为 x, y 的一个齐次多项式(如 $p(x, y) = x^2 - 3y^2$)， m 为任一整数，以 $N_p(m)$ 表满足方程式

$$p(x, y) = m$$

的所有整数解的个数。

我们已知的是 $N_p(m)$ 总为有限的（其值当与 p 及 m 有关），但如何求出那些整数解及精确判断 $N_p(m)$ 的界限；这些一直都是数论中漂亮而且艰深的研究课题。

参考文献

1. S. C. Lu & L. N. Lee, A Simple and effective public-key cryptosystem, COMSAT TECHNICAL RENEW vol. 9, No. 1, 1979, pp. 15—24.
2. B. S. Adiga & P. Shankar, Modified Lu-Lee Cryptosystem, Electronics Letters, 1985
3. L. M. Adleman & R. L. Rivest, How to break the Lu-Lee(COMSAT)public-key cryptosystem, MIT Laboratory for Computer Science, July 1979.
4. M. J. Kochonski, Remarks on Lu & Lee's Proposal, Cryptologia, 1980.
5. R. Solovay & V. Strassen, A fast Monte-Carlo Test for primality, SIAM Jour, on Computer, March 1977, pp. 84-85.

第四章 未来数学家的挑战——计算量

§1 前言

有数学家说过“一个好的问题胜过十个好解答”。因为解答一出，此问题已是到了终点，对不断求创新的人们而言，已不构成挑战。而新的问题是源头活水，能开拓新的境界。多数人都宁愿沉醉在好的解答中不断的玩味，而希望找到新的问题，不断的思考、摸索。

大家在《数学传播》上已看见了不少好的问题，尤其最近康明昌教授谈到的费马定理，几何三大难题，都是极有趣的问题。有的已有了解答，有的尚待解决。除了上面的题目外，像四色问题（即任何一个地图只要用四种颜色就可以把国界分开），五次以上方程式的公式解，及数论上质数分布问题，都曾在职业及业余数学家的心目中占有相当的地位。本文所要介绍的是一个最近（1970年开始）一种许多数学家及电子计算机学家所关心的大问题——NP问题。NP 所代表的意思，你看完本文之后自然

会明白，现在你不妨记住“NP-hard”这个伟大的字。将来如果你对某人说你的问题是“NP-hard”，他也许就要对你刮目相看了，NP-hard不但表示hard(难)，而且是NP的难！

NP问题的代表问题之一是售货员旅行问题(traveling salesman problem)。有一个售货员要开汽车到 n 个指定的城市去推销货物，他必须经过全部的 n 个城。现在他有一个有此 n 城的地图及各城之间的公路距离，试问他应如何取最短的行程从家中出发再回到家中？

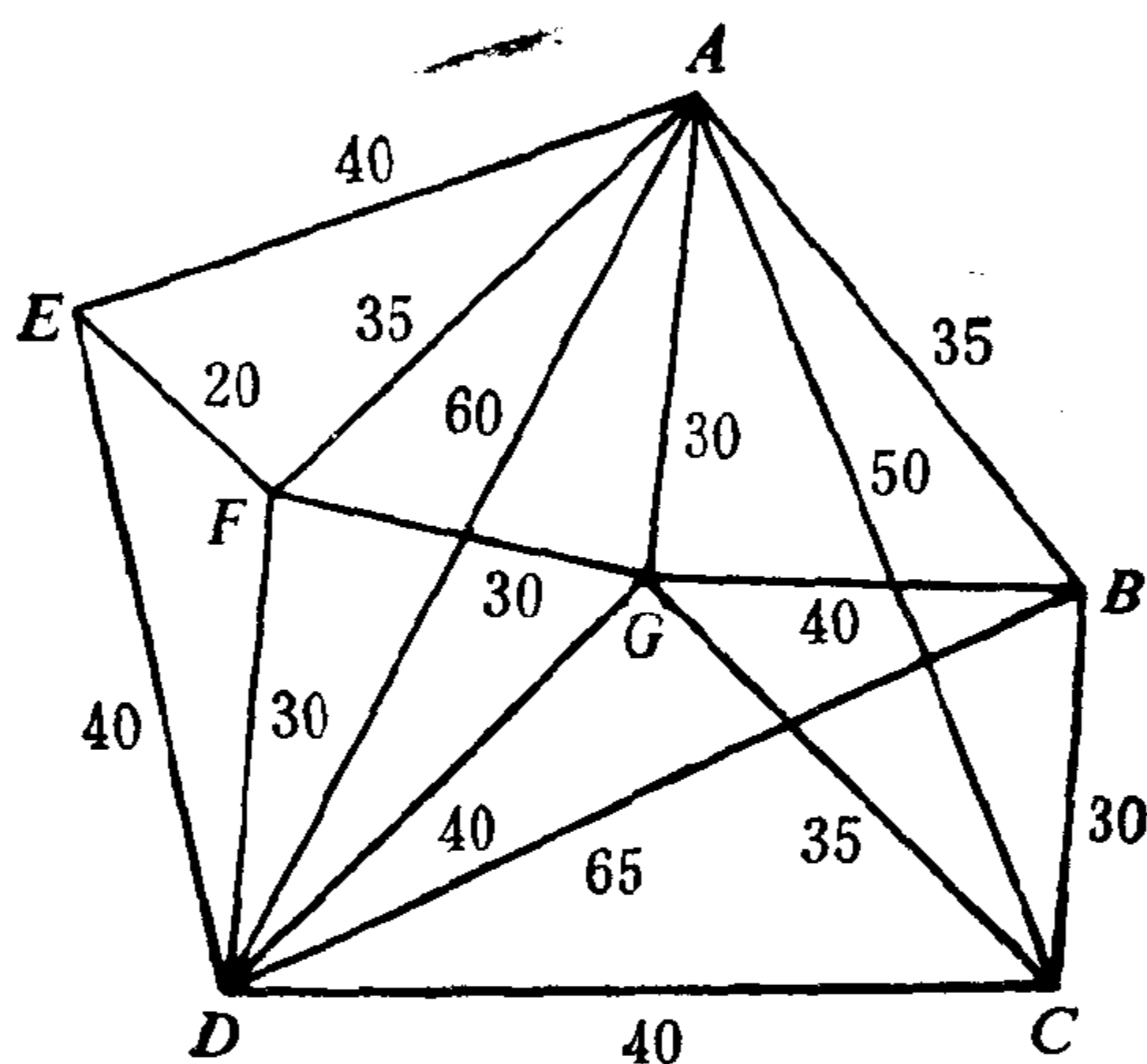


图7 售货员行走的地图A, B, C, ...表城市，数字表两城之间的里数

如图7中，A, B, C, ..., G表示7个城市，而售货员要从A城出发再回到A城并访问B, C, ..., G所有的城，一个可行的方法是

$$A \rightarrow B \rightarrow C \rightarrow G \rightarrow D \rightarrow E \rightarrow F \rightarrow A$$

问题是：这是否是最短的途径？也许

$$A \rightarrow G \rightarrow B \rightarrow C \rightarrow D \rightarrow F \rightarrow E \rightarrow A$$

更近呢？加起来的结果第一路径总长235里，而第二路径总长为

230里，故第二路径较短，但是否存在一个更短的路径呢？目前的方法接近一个一个的排着试，还没有找到更好可以寻得最短路径的方法。对七个城而言，共有 $6! = 720$ 个排法，尚不算难，但若有20个城，则排法就有 $19!$ 种。因

$$n! \approx \sqrt{2\pi n} (n/e)^n$$

故 $19! \approx 1.21 \times 10^{17}$

在排列组合里 $n!$ 写起来轻松，但 1.21×10^{17} 是一个大得不得了数字，若每秒钟排一次，要排 3.84×10^9 年（一年约为 3.15×10^7 秒），即使使用计算机，每秒排一百万次（不容易做到）也得重做三千年才能找到答案。“生也有涯，知也无涯”，想不到区区二十个城，要三十个世纪才能找到答案。

由于电子计算机的发展，有许多以前认为枉费时的计算，像行列式之值、反矩阵、高次方程式的解，都可以在极短的时间内解决。但也突然出现了一些新问题，连大型计算机也望之兴叹。像售货员问题，因为找不到比硬排好得很多的做法，使得数学家们开始想要证明，根本找不到比硬排好得很多的做法。这个证明至今尚未找到。就像以前一角三等分问题一样，既然找了几千年找不到用圆规直尺三等分任一角的方法，也许我们可以证明绝对不可能用圆规直尺三等分一角。现在我们要证明绝不可能写一个计算机程式大大的简化售货员旅行问题。与三等分一角问题不同的是，前者是一种数学上的好奇，而当今的问题与实际用途却有密切的关联。

在此我们一直强调一个好得很多的方法，原因是对这类的问题，你若能计算快一倍或十倍、千倍，往往起不了什么大作用，好像刚才的二十城旅行问题，即使快了千倍，仍需三年的计算时间，而再加三城立刻就把这个计算法的效果抵消了，因此我们所要的是计算量基本层次的减少，这就是我们在下一节

所要讨论的。

§ 2 计 算 量

计算量，顾名思义，是指解决某问题所需要计算的时间，但因每个复杂问题的计算往往都要经过许多不同的运算，除加减乘除四则外，还要包含比较、取数据、存数据等等，若仔细计算起来，十分困难，一般都只绘出一两个主要的量，加以统计，以上节中售货员旅行问题为例，其主要的工作是对每一个排法加起总路径之长，因对 n 城而言，有 $(n-1)!$ 的排法，我们就定其计算量为 $O(n!)$ ，即在 $n!$ 之层次(order 即 O 缩写之来源)之内。

举二个例子，我们若要求 n 个数的和或平均值，则其计算量为 $O(n)$ ，但若我们要把 n 个数字依次排列，则其计算量会因做法的不同而有相当的差别，一个直接了当的方法是，先求出最大的(比 $(n-1)$ 次)，再从不是最大的中间求次大的(比 $(n-2)$ 次)，再求第三大的(比 $(n-3)$ 次)，…如此一共比了

$$(n-1) + (n-2) + \cdots + 1 = \frac{n(n-1)}{2}$$

次就可以完成此工作。因此我们以 $O(n^2)$ ，即在 n^2 之层次来表此方法的计算量。另外一种快排法，先把 n 个数分成若干小块，每块排好之后再合起来，则可以证明此种方法之计算量为 $O(n \log_2 n)^{(1)}$ 。因排数字与排名字、电话号码相同，这种排法很有实用价值，例如某大城有一百万户，则 $n^2 = 10^{12}$ ，而 $n \log_2 n$ 只有 2×10^7 ，其差别三个月与一分钟之比。

表1 以计算机每秒做一百万次时完成各层次计算量所约需的时间
(若无单位, 均以秒为单位)

n	$\log n$	n	$n \log n$	n^2	n^3	2^n	3^n	$n!$
10	10^{-6}	10^{-5}	10^{-5}	10^{-4}	10^{-3}	10^{-8}	0.059	0.45
20	10^{-6}	10^{-5}	10^{-5}	10^{-4}	10^{-2}	1(秒)	58(分)	1年
50	10^{-5}	10^{-4}	10^{-4}	0.0025	0.125	36年	2×10^{10} 年	10^{57} 年
1000	10^{-5}	10^{-3}	10^{-3}	1	16小时	10^{888} 年	极	大
10^8	10^{-5}	1	6	1月	10^5 年	极	大	
10^9	10^{-5}	16小时	6天	3年	3×10^9 年	极	大	

一般计算量的层次多以下表来区分:

$$O(\log n) < O(n) < O(n \log n) < O(n^2) < O(n^k) \\ < O(2^n) < O(k^n) < O(n!)$$

在上表中, k 为某一大于2的正整数, 它们中间都有一道鸿沟, 有基本层次的不同, 在计算机理论上, 若某人能发现一个新的方法, 降低一个层次的计算量, 那么他的新方法有资格称之为一个突破, 可以不朽矣。表1有一个对上项各量的比较, 是以计算机每秒作一百万次(10^6)计算为原则。

在这个表中, 特别注意 n^3 与 2^n 中之差异, 一般称 2^n 为计算量呈指数上升, 而 n^3 或 n^k 之计算量呈 n 的方次上升⁽²⁾, 对目前及未来的计算机而言, 一个呈方次上升的计算量应可以应付, 但对一个呈指数上升的计算量在 n 相当大时则毫无希望。因此计算机学家所集中精力的方向在如何将一个呈指数上升的计算量问题, 简化成一个方次上升的计算量问题。我们对定义凡对一个问题中最重要的参数 n 而言, 若能找到一个方法可以以方次上升的计算量完成, 我们称此问题为一 P -问题(P 为英文多项式 Polynomial之第一字母), 包含所有此问题之集合以 P 表示之。

§3 P 之外?

本节的题目有点不平常，我们的目的是提醒读者本文中常用的英文大写的 P 是一个凡能用 $O(n^k)$ 计算量解决的问题的集合。而 P 之外加一个问号系指到目前为止，我们尚不知道 P 之外是否是一个空集合。

到目前为止，除了售货员旅行问题之外，已经有上百有趣或有用的问题，无法用 $O(n^k)$ 的计算量来解决，我们在此列举几个例子。

问题1：售货员旅行问题(甲)，即第一节所述的问题，不再重复，不过假定所有距离均为正整数⁽³⁾。

问题2：售货员旅行问题(乙)，与第一题的条件相同，但现在有一个给定的正整数 B ，问题是是否存在一条路径其总距离不大于 B 。(问题1与问题2在表面上相似，但在以后的理论上有很大的不同。)

问题3：背袋问题(甲)，有物体 n 个，各重 $w_1, w_2, \dots, w_n$ ，今欲将它们分为二袋，试问如何分法可使两袋的重量最为接近。(不妨假定 w_i 皆为正整数，这并未失去一般性。)

问题4：背袋问题(乙)，如上题，并给定一正整数 B ，试问可否选出若干 w_i ，使其和

$$S \leq \sum_{i=1}^n w_i / 2 \quad \text{且} \quad S \geq B$$

问题5：包装问题：有 n 个各别重量小于1公斤的物品及足够可以装1公斤东西的盒子，今将物品装于盒子之中，各个物品可装于一盒，但任何一盒不得重于1公斤，试求最小的盒子数。

问题6: 舞伴问题: 今有 n 个男孩子与 n 个女孩子参加舞会, 每个男孩与女孩均交给主持一个名单, 写上他(她)中意的舞伴(至少一人, 但可以多于一人)。试问主持人在收到名单后, 是否可以分成 n 对, 使每人均得到他(她)所喜欢的舞伴。

问题7: 库存问题: 某仓库有 D 个存仓, 排成一行, 今有 n 批货物, 各可占有一个或多个存仓, 并已知各批物品存入与提出的日期。试问可否将各货物存入库里不发生存仓不够的困难且同一批货物若需一个以上存仓时, 其存仓必须相邻。

问题8: 已知 a, b, n 三正整数, 问是否存在一小于 n 位之正整数 x , 使得

$$x^2 \equiv a \pmod{b}$$

问题9: (甲): 给定一 n 位正整数 a , 试问其是否为质数?

(乙): 给定一 n 位正整数 a , 试问是否存在 $m, n > 1$ 且 $a = mn$?

问题10: 分丛问题: 已知空间 n 个点, 并假定各点之间的距离为正整数, 又给定两正整数 K 与 B , 问是否可将此 n 点分成小于 K 个不重合的子集, 使得在同一子集内的任意二点距离均不大于 B ?

现在可以看出这类问题的一般结构了。很显然的, 有些是极有用的问题, 而有些可以转换成有用的问题。例如舞伴问题, 若把男孩与女孩换成工人与工头, 或医生与病人就有大用了。这些问题到目前没有一个可以证明是属于 P 的, 大家都猜测它们可能在 P 之外, 即其计算量是呈指数增加的。

在60年代, 已有些人把某些问题归于一类了, 即是几个问题是互依的, 若其中之一属于 P , 则其它几个也属于 P , 其证明方法大都是证明两个互依问题中间有一个只需要用 $O(n^k)$ 时间来完成的桥梁。直到1971年古克 (Stephen A. Cook) 发表了

“The Complexity of Theorem Proving Procedures”才把 P 之外的问题归成了三大类，即NP，NP-complete及NP-hard⁽⁴⁾，现在谈古克定律。

§ 4 古克定律与NP-completeness

古克定律的证明很难，就是了解它也不容易，我们将从几个角度来看这个问题，试着去了解它。它的主要结果是把前节那类问题大都归于一个较易证明的集合，称之为NP，而在NP中找到一批互依的问题称之为NP-complete类并得到下面的结果。

1. 若有一个NP-complete问题可以用 $O(n^k)$ 计算量来解决，则全体的NP问题都可以用 $O(n^k)$ 的计算量来解决，即

1' 若有一个 $x \in \text{NP-complete}$ 且 $x \in P$ ，则 $P = \text{NP}$ 。

又换句话说，NP-complete是NP中的难题，NP-complete解决了⁽⁵⁾，NP就解决了。但若有一个属于NP而不属于NP-complete的问题解决了，则其它的NP问题不一定可以解决。

什么叫做NP？NP是英文nondeterministic polynomial的缩写，意思就是非确定性的多项式。要了解这个字，我们先看一看普通计算机的作用。

现在已知用一个计算机，要解决售货员旅行问题非常困难，但若我们有许多计算机同时用，是否可以快到把原问题在 $O(n^k)$ 时间内解决？“许多”，不是一、二，多一二个是于事无补的，多百个千个仍是杯水车薪，不能有很大的作用，因为就是一千个机子可以分开做，也最多只能快一千倍，在第一节内已说过，帮助不大。因此计算机学家先放眼望去，干脆允许你可以无限

的增加机器。现在我们要注意的是并不是有了无限多的机器，所有的问题就可以立刻解决了，因有的问题有先后次序，例如在算下式的时候

$$[(a_1 + a_2)a_3 + a_4]a_5 + a_6$$

除非换个形式，否则必须一步一步的解括弧，机器多了并不能加快计算的速度，而且机器多了，其间的联络千变万化，一个机器要应付千千万万别的机器送来的信号也疲于奔命了。因此我们只假定所有的机器都只承上启下，单线作业，不作任何横向联络⁽⁶⁾，也就是说，机器1可以把它的结果传给它下面的机器，像 $a_1, a_2, \dots, a_n$ ，而每一个机器又可以把它们的结果传给自己的子机，但在 $a_1, a_2, \dots, a_n$ 之间不互相联络。以售货员旅行问题为例，若有20个城，第一个机器开始，叫下面19个机器各取一个不同的城及计算与A距离，而这个19个机器又将它所求得距离交给自己的18个子机，令它们取一个与自己不同的城加上距离，如此往下，在第十次时，第十阶段的机器把它已取9城及总距离告诉下一个机器，叫它们再取一与已取之城不同之城加上距离，如此一直做到第19次，所有路线的距离都有了，在时间上求得所有的距离是 $O(n)$ （但用了 $19!$ 个计算机），古克定义可以在 $O(n^k)$ 时间内用无限多计算机解决的问题为一NP问题。

现在要记住的是由于无横向连络，在所有路径的距离都有了之后，并没有解决售货员问题(甲)，因为不知谁是最短（若加以比较以求最短距离，则要 $O(n!)$ 个比较），因此我们不能说售货员旅行问题(甲)是一个NP问题。但上节问题2，售货员旅行问题(乙)，任何一个单线都可以知道它的总距离是否不大于B，因此每单线都有一个“Yes”或“No”的答案。只要有一个“Yes”的答案，我们即知道本问题已解决，故问题二是一个NP

问题。在单线作业中，每个机子可以作三件事。

1. 目前答案不明确，大家各自作业。
2. 某线已找到答案，立刻叫停，大家停止作业，解题完毕。
3. 此路不通，本线不再作业，但不叫停，别的线仍然作业。

从上项作用，很容易看出找出答案的计算时间即某线叫停的时间，亦即任何一个有“Yes”答案线中计算量之总和，也就是说找到答案“捷径”上所需的时间。说得更明白一些，在一非确定性计算机系统下，其子机像有“猜测”到捷径的功能。若在任何计算步骤中，某人猜了一个答案，而计算机可以在 $O(n^k)$ 时间内回答“Yes”或“No”，这个问题即是一个NP问题。再以售货员旅行(甲)及图7为例，若你猜一个路径

$$A \rightarrow B \rightarrow C \rightarrow D \rightarrow E \rightarrow F \rightarrow G \rightarrow A$$

我们无法知道此路是否最短，但在(B)问题中，一个“Yes”或“No”的结果只要7个加法就可以回答了。因此根据新的定义，问题2是一个NP问题。

由这两个定义，读者不难看出问题2, 4, 6, 7, 8, 9(乙)与10皆为NP问题，特别是问题9的(甲)、(乙)，其实是一样的问题，但如果你猜二个 m, n ，立刻就可知道 a 是否是 mn 。

古克定理的关键在证明若一种叫满足问题(Satisfiability Problem)的例子属于P，则所有NP问题均属于P（即此问题属于NP-complete），令 $\cdot, +, -$ （且，或，反）表三个基本的逻辑运算（即对0与1逻辑符号而言， $\bar{0}=1, \bar{1}=0$ ，除了 $1+1=1$ 之外， $+, \cdot$ 与一般代数之加乘相同）。令 f 为一个含有 n 个逻辑变数 $(u_1, u_2, \dots, u_n)$ 的函数。假如我们可以找到一组 $u_1, u_2, \dots, u_n$ ，使得 $f(u_1, u_2, \dots, u_n)=1$ ，则 $f(u_1, u_2, \dots, u_n)$ 称为可满足。例如

$$f(u_1, u_2, u_3) = (u_1 \cdot (\bar{u}_1 + u_2) \cdot u_2 + u_3) \cdot u_3 \quad (1)$$

为一可满足函数，取 $u_1 = u_2 = u_3 = 1$ 即可，但

$$f(u_1, u_2, u_3) = (\bar{u}_1 \cdot u_1 + u_2 \cdot \bar{u}_2)(u_1 + u_2 + u_3) \quad (2)$$

永为0，故此 f 为不可满足。

直觉上，这类问题除了将 $u_1, u_2, \dots, u_n$ 一个一个以0, 1代入检查 2^n 次之外，显无捷径可循，古克1971年的论文即证明这是一切NP问题之母。

古克定理：满足问题为NP-complete。

现在已可证明在前节中的问题，除了问题1, 3, 5, 9之外，全是NP-complete问题。

§5 NP问题之近似解

NP问题既找不到可行的解法，而很大部分的NP问题都在计算机语言、程式、电路设计、统计学、程式作业上有大用，因此只好退而求其次找一个可行的近似解。很可惜的是，所有的NP-complete问题虽在NP的层次上相联，在近似解上往往各需不同的解法，这些解法多从直观而来，我们在此举二个例子。

例1 在第三节问题5，包装问题中，若采取“能装就装”法，即现有的盒子若可以装得下，就不用新盒子，则此法所需用之盒子数 k_1 与最可能少的盒子数 k_0 满足 $k_1 \leq 2k_0 + 1$ 。

证明：今令 n 个物品之重为 $w_1, w_2, \dots, w_n$ 公斤，因每个盒子只可以装1公斤，故

$$k_0 \geq \sum_{i=1}^n w_i$$

另一方面，“能装就装”法不可能有两个以上的盒子同时少于1/2公斤，故

$$k_1 \leq 2 \sum_{i=1}^n w_i + 1$$

本例得证。

这个问题的结果是说，我们大约可以用“能装就装”法做得最好情形的一半好。经过较复杂的证明，Johnson 在1974年证得，当 n 很大时，

$$(1) \quad k_1 \leq \frac{17}{10}k_0 + 2, \text{ 且存在一种情形能产生。}$$

$$(2) \quad k_1 \geq \frac{17}{10}(k_0 - 1).$$

也就是用“能装就装”法不会坏到70%以上，但可以坏到多用了70%的盒子。

例2 售货员旅行问题的一个直观走法是先访问那个尚未访问过最近的城，称为“先访近城”法，以图7为例，其走法为

$$A \rightarrow G \rightarrow C \rightarrow B \rightarrow D \rightarrow F \rightarrow E \rightarrow A$$

Rosenkrantz等在1977年证明这并不是一个很理想的走法，他们证出若各城间的距离满足三角不等式，⁽⁷⁾则“先访近城”法所走之总程 D_1 与最短路径 D_0 之关系为

$$D_1 \leq \frac{1}{2}([\log_2 n] + 1)D_0$$

且当 n 很大时，可以有一种情形使得

$$D_1 \geq \frac{1}{3}(\log_2 n + \frac{4}{3})D_0$$

上式中之 $[x]$ 表示大于 x 之最小整数，假如 $[5]=5$ ， $[2.5]=3$ 。因 $\log_2 n$ 当 n 大时可以很大，故 D_1 可与 D_0 相差非常之大，但在同一篇论文之中，Rosenkrantz等证明另一种复杂的“直观”走法可以达到 $D_1 \leq 2D_0$ 的地步。

在上面的定理中，三角不等式的条件很重要，若城之距离无此关系存在时，Sahni与Gonzalez在1976年证得：若 $P \neq NP$ ，则不可能存在一个有限的 m ，及一个 $O(n^k)$ 计算量的走法，能使其全程长 D_1 在任何 n 时满足

$$D_1 \leq mD_0$$

即上式中 m 非等于无限大不可，亦即所有 $O(n^k)$ 的做法都不很好。

§6 NP-hardness与围棋

不是所有的难题都可归结为NP问题，像下得一手绝对好的围棋现在目前的推测是比所有NP问题还要难的计算问题，即NP-hard问题，NP-hard问题的定义如下：

若 x 为一NP-hard问题，则若

$$NP \neq P, \text{ 则 } x \notin P$$

也就是说，即使 $P = NP$ ， x 还不一定属于 P ，但 $P \neq NP$ ，则 x 绝不比NP的问题容易。在第三节中的问题1、3不一定是NP问题，但若能以 $O(n^k)$ 的计算量解决它们，则比较容易的问题2与4也可以 $O(n^k)$ 解决，故若问题1，3 $\in P$ ，则问题2，4 $\in P$ ，又因2，4是NP-complete，即推出 $NP = P$ 。这与NP-hard之定义相合，故问题1，3均为NP-hard问题。同理问题5也属于NP-hard，不过这些NP-hard似乎比NP难不了多少，但下棋问题可能比NP问题要难得多，围棋问题可以作如下叙述。

问题11：（围棋问题） 以平常的围棋规则在一个 $n \times n$ 的棋盘上下，给定一个残局（下了二个子就可以算残局），黑先，是否可以确定黑子在最好的下法之下，一定会赢？

这个问题不能用一般的方法证明它是不是为 NP. 因为目前没有人能猜一个必胜的下法且在 $O(n^k)$ 时间内证明它是对的, 因为它与对方如何应付有关, 而敌方的应付又与他对你以后的下法的推测有关, 如此往下走, 首先发生困难的是记忆上亮了红灯, 即所需要的记忆可能呈方次以上的进展.

因每一个记忆至少要用(来计算)一次, 否则这个记忆就不如不要, 因此一个问题的记忆若呈指数上升, 则其计算量亦非呈指数似的上升不可.

因此计算机学家定义三个新的集合:

$PSPACE = \{x: x \text{ 只需要方次上升的记忆} \}^{(8)}$

$PSPACE\text{-complete};$

若 $x \in PSPACE,$

又 $x \in PSPACE\text{-complete}$

且 $x \in P,$

则 $P = PSPACE.$

$PSPACE\text{-hard};$

若 $x \in PSPACE\text{-hard}$

且 $x \in P$

则 $P = PSPACE$

注意在上式中 $PSPACE\text{-complete} \subset PSPACE$, 即 $PSPACE\text{-complete}$ 是 $PSPACE$ 中的难题, 但 $PSPACE\text{-hard}$ 不一定属于 $PSPACE$. Stockmeyer and Meyer 在 1973 年证明了一个与古克相似的定理.

若令 $\exists x$ 表示存在一个 x , $\forall x$ 表对所有的 x , Q 表 $\exists, \forall$ 中的一个, x 为布氏变数 0 与 1, 则我们称 $f(Q_1x_1, Q_2x_2, \dots, Q_nx_n)$ 为一量化布氏公式. 若 f 有可能为 1, 则 f 称之为可满足, 例如把第四节中之(1)式改写成

$$f(\forall u_1, \exists u_2, \forall u_3) \\ = ((\forall u_1) \cdot (\overline{\forall u_1} + \exists u_2) \cdot (\exists u_2) + \forall u_3) \cdot (\forall u_3))$$

则上式不可能满足，因对 $\forall u_3$ (u_3 为0或1)而言， f 不全是1。

Stockmeyer与Meyer之定理为：

定理 检定一个量化布氏公式为可满足是一个PSPACE-complete问题。

当我们下棋面对着一盘残局沉思的时候，我们要求是

对我是否存在一着必胜棋可以对付

敌人任何一着应付棋

此后我是否存在一着必胜棋可以对付

敌人任何一着应付棋

.....

我是否存在一着必胜棋可以对付

敌人任何一着棋

我赢了

因此这完全是 $\exists, \forall, \exists, \forall, \dots$ 之交替作用与Stockmeyer与Meyer定理之关系至为密切，Robertson与Munro在1978年证得围棋是一种PSPACE-hard的问题，目前有人计算到围棋⁽⁹⁾必胜法的记忆计算量在 10^{600} 以上，不论人脑或电脑的记忆绝少不了一个原子，而现今所知的宇宙原子数约只有 10^{75} 个。棋之道，大矣哉！要做一个下围棋必胜的机器人谈何容易！

§7 结 论

现在你明白20世纪的大难题了， $P=NP$ ？用简单的语言说，就是是否能找到一个只呈方次增加的计算方法去解决旅行、

包装、舞会等问题。平凡的问题，期待您不平凡的解答。

小注

1. 依次排 (Sorting) 的方法很多，但都不能低于 $O(n \log n)$ ，读者可在一般Database的书或第七章中找到有关Sorting的法则。
2. 又称呈多项式上升，但因一个 n 的多项式的大小，在 n 很大时都为第一项所支配，故可写成 $O(n^k)$ 。
3. 并不失去一般性，即若距离不是正整数也可以把它们化成正整数。
4. 在古克的原文中，并没有NP-complete, NP-hard 之明确定义。但是由于他的论文，使这种分法显得很自然。不过NP-hard之定义仍因人而异，不一定同于本文。
5. 本文中之解决，均指一个 $O(n^x)$ 计算量的解法。
6. 与情报人员的单线作用相同，一个谍报人员只知道他的顶头上司及他的第一线下层下属，其余的人他都不应该知道。
7. A, B, C 表任三城，而 $d(A, B), d(B, C), d(A, C)$ 分别表示 $A, B; C; C, A$ 城之距离，则 $d(A, B) \leq d(A, C) + d(B, C)$ 称为三角不等式。
8. x 均指问题。
9. 指 19×19 的棋盘，许多计算机学家都是围棋高手，中国的算盘与围棋，好像包含了计算机的开始与终极。

参考资料与引用论文

1. Gorey, M. R. and Johnson, D. S. "Computers and Intractability——A Guide to Theory of NP-completeness",

1979, Freeman and Company.

2. Pearl, J. "Heuristics—Intelligent Search Strategies for Computer Problem Solving", 1984, Addison-Wesley.
3. Cook, S. A. "The Complexity of theorem-proving procedure" Proc. 3rd Ann. ACM Symp. On Theory of Computing, 1971, 151~158.
4. Johnson, D. S. et. al. "Worst case performance bounds for simple one-dimensional packing algorithms" , SIAM J. Comp. 1974, 299~325.
5. Rosenkrantz, D. J. et. al. "An analysis of several heuristics for the traveling salesman problem," SIAM J. Comp. 1977, 563~581.
6. Sahin, S. and Gonzalez, "P—Complete approximation Problems" J. , ACM 1976, 555~565.
7. Stockmeyer, L. J. and Meyer, P. R. "Word problems requiring exponential time, " Proc. 5th. Ann. ACM Symp.
8. Robertson, E. and Munro, I. "NP-completeness, puzzles, and games" Utilitas Math. 1978, 99~116.

第五章 自动校正码理论浅介

§1 前 言

近代资讯科学的快速发展，得力于电子计算机的普及与多频道通讯网的形成，二者相辅相成，缺一不可。若少了快速通讯网，则每个电子计算机必须由键盘、磁带或磁碟输入资料，大大地减低了快速的功能，但是若没有计算机在通讯网的两端，则快速的通信无法及时消化，亦不能发挥其最大的功能。可惜目前尚未能完全保证通讯中不发生错误。发生错误的原因很多，有自然因素如雷电，有人为因素如强电流及别的通讯信号所带来的干扰。本文所要讨论的就是如何发现或自动更正信号在通讯中所发生的错误。

电子计算机加上通讯网，其所处理的资料之多，信号传递之快，已不是人肉眼校对可以胜任得了的，我们必须在计算机通讯的本身加上检错的装置。一个浅而易见的检错法就是把要传送的信号重复一次，如两者不合则必定有错。比如今有某个

被重复的信号是某人存了一万元及某人存了一千元，二者不符，至少有一个错，也许你会说，如果两次传递都发生了同样的错，把一万误传成了一千，不就查不出来了吗？不错，但且看两次同时出错的可能性是多少？假定一次出错的或然率是 p ，则二次同时出错的或然率是 p^2 (1)。这儿的 p 必须是很小的数量，否则信号老是出错，改不胜改，则此种机器尚不宜问市或要修理。设 p 是十万分之一，即 10^{-5} ，则 $p^2 = 10^{-10}$ 即十亿分之一的或然率，碰到了这种同时出错的事，真是天亡我也。

有时知道有错还不够，因一个信号被送两次而不符合，必是错了，但却不知谁对谁错，只好请对方重传一次，这往往是一个耗时费事的程序，发方必须打断原有的传递，重新找到对方所说传错的位置，重新再传一次。若是“知错能改”，岂不是“善莫大焉”的事吗？何不每个信号传三次，取其多数？由或然率看来，每个信号传三次，是极具有自动更正的功能，但代价是否太大？一个信号传三次，也许太浪费了吧？有没有更经济的方法？

为了使我们的回答更合于电子计算机通讯，我们就采用它们之间通信的符号，即以数码0，1所组成的数字来代表信号，像现在电脑中英文通讯所用的ASCII电码是由七位数码组成的。例如 a 是 1100001， b 是 1100010，1 是 0110001。因每位数有两个可能性，所以一共可以排出 $2^7 = 128$ 个不同的电码，足够表示英文大小写字母、标点符号、数字、希腊字母了。为了说明简洁起见，我们今用两个数码来表示4个信号，即如下页的表1。

现在看，若是 a 为 00 误传成了 01，我们会以为信号是 b ，无法确定会有错误，若我们每个信号传二次，即 a 以 0000 传， b 以 0101 传，等等。则如果 a 的 0000 错传成了 0100，则我们就知道出毛病了，但却不知它是由 0000 或者 0101 而来，由于二者

表1

信号	表示数码
a	00
b	01
c	10
d	11

都只含一个错误。当然0100也可能是由1111误传而来，但因一个错误比多个错误的机会小得多，我们最好的选择是假定只有一个错误，这是引用了统计学上常用的最大可能原则(maximum likelihood principle)。它也是我们日常推测事情真相的原则。比如我们在地上拾到了一笔钱，我们总会推测这是有人遗失的(要交给警察)，而不会认为是有人故意放一笔钱在地上做好事送给人(可以谢谢他而据为己有)。因为后者的可能性太小，在没有其它的证据之下，我们不会做后者的推测。

再仔细研究一下，我们之所以知道0100信号有错，是因为0100不在对方输出的码中，因对方只可以输出0000，0101，1010，及1111的其中之一，可见若要能检错，输出码不可以占用了全部可能的码字，为了统一起见，我们作以下的定义：

定义5.1 凡是用 n 位0或1作为一个单元(或数码)输出入信号者，称为一 n 位码(输送)系统(或简称为码)。其中任何一个信号，若含有 n 个数码，称之为一个码字(注意文中有时将“字”省去)。

定义5.2 在 n 位码系统中，其前 $m(m \leq n)$ 位称之为信息码(或原码)，它是我们所有要传的信号，而后 $t = n - m$ 位称为检定码，是用来检定或改正错误而用。

下文所谓检定 k 位是指说可以知道至多有 k 位数码错了，而校正 k 位是指说把至多有 k 位错的数码的位置找出来。

以表 1 及重复传送为例，其前两位是信息码，后两位是检
定码(见表2)，其中 $n = 4$ ， $m = 2$ ， $t = 2$ 。

表2

信 号	信 息 码	输 送 码	检 定 码
a	00	0000	00
b	01	0101	01
c	10	1010	10
d	11	1111	11

用表2的输送码，我们已知他们可以检错一位，但不能更正
错误，对检错一位的要求而言，我们又不知道这种重复输送是
否是最经济的办法（答案是否定的，我们在第三节会谈到），
现在，我们可以把我们想解决的命题写下来了：

对一个原为 m 位的信息码，我们最少要加上多少位的检定
码 t ，使之可以自动检定或更正 k 位的错误？ (1)

在上面的命题中， m 与 k 为已知数，而 t 为待定的数。

要解决这个问题，我们首先把这些码字用向量的符号表示，
并定义两码字的距离，然后我们证明当可以输出的码字之间都
有相当大的距离时，某些位的错误就可以检定或校正。

看本文所需要的预备知识是矩阵的一些基本运算，最后在
结论中，我们提出了一些较难的问题，供有兴趣的读者思考或
做更深入的研究。

§ 2 如何表示码字与它们之间的关系

我们延用定义 1、2 中的符号 m ， n ， t ，如果以 Z 表 $\{0,1\}$ 之

集合, 则 m 位信息码可以 Z^m 表示之, 即

$$Z^m = \{(x_1, x_2, \dots, x_m), x_i = 0 \text{ 或 } 1\} \quad (2)$$

其中一共有 2^m 个元素. 同样我们用 Z^n 表示所有 n 位码字之集合, 则制码法就是把一个 Z^m 中的元素镶到 Z^n 中去的函数, 以表2为例, 重传就是把 Z^2 中的元素 (a_1, a_2) 镶成 Z^4 中的元素 (a_1, a_2, a_1, a_2) , 以一般函数符号表示, 则

$$f = Z^m \rightarrow Z^n \quad (3)$$

表示一种制码法, 令 A 为 f 之区域, 即

$$A = f(Z^m) = \{f(x) | x \in Z^m\}$$

则很显然的 $A \subset Z^n$, 我们又称 A 为输出码表, 而且在我们表示向量 $(x_1, x_2, \dots, x_m)$, 有时不加逗号, 即在表2中的 b , 我们可以写成 $(0, 1, 0, 1)$ 或 (0101) .

定义5.3 在(3)中的制码法如果在收到的 n 位码字中有 k 位或 k 位以下的错误时, 能知道收到的码字有错, 则此种制码法称之为有检定 k 位错误的功能, 若其能校正此 k 位错误, 则称之为校正 k 位错误之功能.

例如表2中的输送码, 有检定一位错误之功能, 但没有校正一位之功能.

定义5.4 (汉明(Hamming)距离) 设 X, Y 为 Z^n 中的两元素, 则 X 与 Y 之间的距离定为 $H(X, Y) = X, Y$ 中不相同的数码之数目.

例如: $X = (1010), Y = (1001)$, 则因 X 与 Y 有二位(第三、四位)不相同, 故 $H(X, Y) = 2$, 若令 $0 = (0, 0, \dots, 0) \in Z^n$, 即全含0的元素, 则

$$H(X, 0) = X \text{ 中所含 } 1 \text{ 的数目}$$

并以 $|X|$ 表示之.

定义5.5 令 $X = (x_1, x_2, \dots, x_n)$

$$Y = (y_1, y_2, \dots, y_n)$$

则定向量之加减:

$$X \pm Y = (x_1 \pm y_1, x_2 \pm y_2, \dots, x_n \pm y_n) \quad (4)$$

且在上式中,

$$1 \pm 1 = 0 \pm 0 = 0,$$

$$1 \pm 0 = 0 \pm 1 = 1.$$

也就是说, 如果 x_i 与 y_i 相同, 则它们的差为0, 否则为1.

由此很容易看出 X 、 Y 的距离

$$H(X, Y) = |X - Y| \quad (5)$$

一些不熟悉抽象数学的读者也许会觉得 $1 + 1 = 0$ 很怪, 但如我们抽象地把0、1分别相应“偶”、“奇”两个概念, 就不会惊讶了. 且我们只有0与1两个数, 因此 $1 + 1$ 只有等于0或1两个选择, 若令 $1 + 1 = 1$, 则会造成 $1 = 0$ 的结果, 二元变成了一元, 一些主要的演算特性都消失了, 因此只好定义 $1 + 1 = 0$, 而熟悉抽象数学的人知道这样的定义加上一般的乘法

$$0 \cdot 1 = 1 \cdot 0 = 0 \cdot 0 = 0, \quad 1 \cdot 1 = 1 \quad (6)$$

使得 $\{0, 1\}$ 成为一个域, 这样上面定的向量集合就自然形成一个向量空间了.

由上面的运算规则, 很容易可以得到及了解下面一些定理及定义.

定理5.1 令 X, Y, Z 为 Z' 中之任意三元素, 则

$$\begin{aligned} (1) \quad H(X, Y) &= |X + Y| \\ &= |X - Y| \\ &= H(Y, X) \end{aligned}$$

$$(2) \quad X + X = X - X = 0$$

注意由此可得 $X = -X$ 的一重要事实.

$$(3) \quad \text{若 } X \neq Y, \text{ 则 } X \pm Y \neq 0$$

$$(4) H(X, Y) \leq H(X, Z) + H(Y, Z)$$

$$(\text{三角不等式}^{(2)}) \quad (7)$$

在结束本节之前，我们定义 Z^n 一个子集的离散度。

定义5.6 令 $A \subset Z^n$ ，则 A 的离散度定为

$$d(A) = \min_{\substack{X \in A \\ Y \in A \\ X \neq Y}} |X - Y| = \min_{\substack{X \in A \\ Y \in A \\ X \neq Y}} |X + Y| \quad (8)$$

则 $d(A)$ 为 A 中不同元素间距离的最小值。

定义5.7 令 $A \subset Z^n$ ，若 $X \in A, Y \in A$ 可以导至 $X + Y \in A$ ，则 A 称为 Z^n 的一个封闭子集。

定理5.2 若 A 为一封闭子集，则

$$d(A) = \min_{\substack{X \in A \\ X \neq 0}} |X|$$

证明很显然，因(8)式中 $X + Y$ 可以 A 中的元素表示之。例如表2中的输送码为一封闭子集，而其离散度为2，再回到表2，表2的制码法可以用矩阵的乘法表示，我们若以 X_m 表二位的信息码，而 X 表四位的输送码，则

$$\begin{aligned} X &= x_m \begin{bmatrix} 1 & 0 & 1 & 0 \\ 0 & 1 & 0 & 1 \end{bmatrix} \\ &= (x_1, x_2) \begin{bmatrix} 1 & 0 & 1 & 0 \\ 0 & 1 & 0 & 1 \end{bmatrix} \end{aligned}$$

当然，一般的制码法(3)并不一定可以以这种乘法表示，但可以用这种乘法表的制码法至少有两个好处，第一是规则简单，第二是它所形成区域 $A = f(Z^m)$ 是 Z^n 中的一个封闭子集，这个特性我们会在第四节中用到。

§3 检 错 码

定理5.3 设 A 为一码表, 则它能检验出至多 k 位错误的充要条件是 $d(A) \geq k+1$, 即 A 的离散度必须大于或等于 $k+1$.

证明 充分性: 设 $d(A) \geq k+1$, $\bar{X}$ 为 X 之误传, 因 $|\bar{X} - X| \leq k$, 而 $d(A) \geq k+1$, 故 $\bar{X} \notin A$, 必定有错.

必要性: 若 $d(A) \leq k$, 则存在 X, Y , $|X - Y| \leq k$. 因我们允许 k 个错误, 故 Y 可以为 X 之误, 但因 $Y \in A$, 我们不知可能的错误, 与假设要求不合.

在这儿要注意一点的, 所谓可以检出 k 个错误是指定义3所说的, 如果 X 有 k 个错, 我们只知道它至多有 k 个数码错了, 但并不知道错在何处. 对数码为0、1的码字而言, 一旦知道错在何处, 就等于可以改正了, 以表2为例, 因 $d(A)=2$, $k=1$, 因此, 如果有一位误传, 我们可以知错(但不知错在何处), 若有两位发生错误, 例如0101传成0000, 则我们就发现不出错误了, 这要 $d(A) \geq 3$ 才行.

由定理5.3可知若要检定输送的码字是否有一个或一个以下的错误, 我们得要有 $d(A)=2$, 若不加检定码, 因 A 的元素皆为信息码, 于是 $d(A)=1$, 不可能检定错误. 因此我们至少要加一位检定数码. 现在我们要证明只要加一位检定数码就足以检定一位的错误, 因此表2中的重复输送是不必要的. 令 $O(Y)$ 表 X 的奇偶数(Parity), 即

$$O(X) = \begin{cases} 1 & \text{如果 } X \text{ 含有奇数个 } 1 \\ 0 & \text{如果 } X \text{ 含有偶数个 } 1 \end{cases} \quad (9)$$

则很容易证明下面的定理.

定理5.4 若 $X \approx Y$, $O(X) = O(Y)$

$$\text{则 } H(X, Y) \geq 2 \quad (10)$$

推论 若 A 中之元素均有相同的奇偶数, 则若在输送时出了一位错, 我们可以检出。

要使得所有码表中的元素都有相同的奇偶数是很容易的事, 只要在信息码后加上一位奇偶检定码使之成为一个含有偶数个1的输送码就可以了⁽²⁾, 因此对表1的四个信号而言, 可以检定一个错的编码法是令 $a=000$, $b=011$, $c=101$, $d=110$ 。

奇偶检定码与原码字长度 m 无关, 前面谈到的 ASCII 码都是七位再加一位奇偶检定码。因此 a 实际是八位的 11000011。

若要使表1中的信息码有检定二个错的功能, 必须要 $d(A) \geq 3$, 字码表(10)是不行的, 因 $d(A)=2$, 再加一位检定码行不行? 答案是不行, 但这不是一眼可以看出来的, 加两位尾数的方法很多, 但我们可以用穷举法证明。若是我们令 a 为 $00a_1a_2$, b 为 $01b_1b_2$, c 为 $10c_1c_2$, 则不可能使得它们之间的距离都不小于3, 加上三位如何? 这就更难验算了, 我们把这个问题与下节的自动校正码合并, 我们将证明可以检二位错的检错码, 就有自动改正一位错误的功能。

在结束本节之前, 我们要顺便提一下, 当我们加码来检定错误; 也会因码位加多而增加了一个码字发生错误的或然率, 但在出错或然率 p 很小时, 仍为“合算”, 读者很容易算出增加码位后发生某种错误的或然率, 因非本文主旨, 我们不再详谈。

§4 校正码

定理5.5 一个输出码表 A 能校正至多 k 位错码的充要条件是 $d(A) \geq 2k + 1$ 。

证明 充分性：设 $d(A) \geq 2k + 1$ 。我们要证明若 $X \neq Y$ ， X 误传成 $\bar{X}$ ， Y 误传成 $\bar{Y}$ ，其错误均不超过 k 个，则 $\bar{X} \neq \bar{Y}$ ，因

$$|X - \bar{X}| \leq k, \quad |Y - \bar{Y}| \leq k$$

由三角不等式及 $|X - Y| \geq 2k + 1$ ，得

$$\begin{aligned} |\bar{X} - \bar{Y}| &\geq |\bar{X} - Y| - |Y - \bar{Y}| \\ &\geq |X - Y| - |X - \bar{X}| - |Y - \bar{Y}| \\ &\geq 1 \end{aligned}$$

故 $\bar{X} \neq \bar{Y}$ ，即 X 与 Y 不可能混合误传。

必要性：设 $d(A) \leq 2k$ ，则可找到 $X, Y \in A$ 且 $|X - Y| = q \leq 2k$ ，我们不妨假定 X 与 Y 在 $a_1 a_2 \cdots a_q$ 位置上不同，取 $\bar{X}$ 与 X 仅在 $a_1 \cdots a_k$ 位置上不同，则

$$\begin{aligned} |\bar{X} - X| &= k \\ |\bar{X} - Y| &= q - k \leq k \end{aligned}$$

故 $\bar{X}$ 之错不能改正，因不知它是由 X 或 Y 而来，本定理证毕。

因此若要编一个可以改正 k 个（或 k 个以下）的校正码，只要在每个为 m 位的信息码后加若干位的数码，使得它的输送码有 $2k + 1$ 的离散度就可以了。

定理 5.6 对具 m 个数码的信息码 A ，若加上 t 位的检错数码，使之可以自动校正 k 或 k 个以下的错误，则 t 必满足。

$$1 + \binom{n}{1} + \binom{n}{2} + \cdots + \binom{n}{t} \leq 2^t \quad (11)$$

上式中 $n = m + t$ 。

证明：在 A 中有 2^m 个码字。任一个码字若错了 $0, 1, \cdots, k$ 位时，必不能与另一个码字错了 $0, 1, \cdots, k$ 位时相同。因每个码字有 $\binom{n}{1}$ 个与它差一位的， $\binom{n}{2}$ 与它差 2 位的， $\cdots$ ，的字码，加上自己，可见每一个 A 中的元素都占有了

$$1 + \binom{n}{1} + \binom{n}{2} + \cdots + \binom{n}{t}$$

个码位，因 Z^n 只有 2^n 个元素，故

$$\left[1 + \binom{n}{1} + \binom{n}{2} + \cdots + \binom{n}{t} \right] 2^m \leq 2^n$$

消去 2^m ，得(11)。

注意：一般书中对原码字为 m 位，检定码字为 t 位的线性二元码的新码以 (m, t) 码表之。

在(11)中，因 $n = m + t$ ，并没有公式可以表示 t 的解，但用代入法，很容易算出在给定 m 时，最小可以满足(11)之 t ，我们称之为 t 之下界，它们的值可在表5中看到。比如说 $m = 10$ ， $k = 2$ ，则至少得加8个检定数码不可，但如何加这8个检定数码，却是一个难题。因为在(3)式中的 f 可以千变万化。在此我们只讨论一种规则的编码法，一般称之为线性码，或群码⁽³⁾，我们要求在编码：

$$f: Z^m \rightarrow Z^n \quad (12)$$

中的 f ，为一矩阵 F_0 ，用符号表示：

$$f(X_m) = X, \quad X_m \in Z^m, \quad x \in Z^n$$

F 为一 $m \times n$ 的矩阵，且

$$X = X_m F \quad (13)$$

为了保持前面 m 位原码的数码不变，我们将 F 写成2个子矩阵， $m \times m$ 的恒等矩阵 I_m 及 $m \times t$ 的矩阵 D ，即

$$F = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 & D \\ \vdots & \vdots & \\ 0 & 1 \end{bmatrix} = [I_m : D]$$

亦即

$$X = X_m F = [X_m : X_m D] \quad (14)$$

如何找到 D ，使得 $d(A) \geq 2k + 1$ 是制造校正码的关键，它的原理却很简单，先定义一个 $n \times t$ 的矩阵 M 如下：

$$M = \begin{bmatrix} D \\ I_t \end{bmatrix} \quad (15)$$

上式中 I_t 为一个 $t \times t$ 的恒等矩阵，则若 $X \in A$ ，由(14)及(7)②可知

$$\begin{aligned} XM &= [X_m : X_m D] \begin{bmatrix} D \\ I \end{bmatrix} \\ &= x_m D + x_m D \\ &= 0 \end{aligned} \quad (16)$$

将 X 与 M 以另一形式展开，即令 $X = (x_1, x_2 \cdots, x_n)$ ，把 M 分解为 n 个 $t \times 1$ 的向量

$$M = \begin{bmatrix} v_1 \\ v_2 \\ \vdots \\ v_n \end{bmatrix}$$

则(16)成为

$$\begin{aligned} XM &= (x_1 \cdots x_n) \begin{bmatrix} v_1 \\ \vdots \\ v_n \end{bmatrix} \\ &= \sum_{i=1}^n x_i v_i = 0 \end{aligned} \quad (17)$$

由这个关系，我们可以证明编校正码的一个基本定理。

定理5.7 在(14)中的 D ，当它写成(15)、(17)式时，若没有任何 $2k$ 个之内的 $v_1, v_2, \cdots, v_n$ 加起来为0，则 $d(A) \geq 2k + 1$ 。

证明 由(17)式可知, $\sum_{i=1}^n x_i v_i = 0$, 但没有 $2k$ 个 $v_1, \dots, v_n$

加起来是0的, 所以 $(x_1, x_2, \dots, x_n)$ 中至少要有 $2k+1$ 个1才行。因此除了 $X=0$ 之外, $|X| \geq 2k+1$, 因 A 为一封闭子集⁽⁴⁾, 由定理5.2可知

$$d(A) = \min_{\substack{X \in A \\ X \neq 0}} |X| \geq 2k+1$$

本定理得证。

定理5.8 若 $k=1$, 则任何 n 个不相同的 v_i 可使得(14)的转换满足 $d(A) \geq 3$, 又用此法时 t 必满足 $m \leq 2^t - t - 1$ 。

证明 因任何两个不相同的 v_i 相加不等于零向量0。由定理5.7很容易证得本定理。因 t 位检定码字只有 $2^t - 1$ 个异于0的向量, 而我们需要 n 个不同的, 故必须有 $n \leq 2^t - 1$, 即 $m \leq 2^t - t - 1$ 。

推论 当 $k=1$ 时, 用(14)所制成的最经济(指 t 大小而言)的校正码, 就是最经济的校正码。

证明 比较定理5.8与定理5.6, 可知当 $k=1$ 时, 不满足 $m \leq 2^t - t - 1$ 的 t , 也就是不满足(11), 即不可能以任何方式做成校正码。

现以 $m=4$ 为例, 由 $m \leq 2^t - t - 1$ 可知最小的 t 是3, M 中 v_i 的选法只有一种, 我们要用全部的 $2^3 - 1 = 7$ 个向量。可令

$$M = \begin{bmatrix} -1 & 1 & 1 & - \\ 1 & 1 & 0 & \\ 1 & 0 & 1 & \\ 0 & 1 & 1 & \\ 1 & 0 & 0 & \\ 0 & 1 & 0 & \\ -0 & 0 & 1 & - \end{bmatrix}$$

因此

$$F = \begin{bmatrix} 1 & 0 & 0 & 0 & 1 & 1 & 1 \\ 0 & 1 & 0 & 0 & 1 & 1 & 0 \\ 0 & 0 & 1 & 0 & 1 & 0 & 1 \\ 0 & 0 & 0 & 1 & 0 & 1 & 1 \end{bmatrix}$$

表3中有这种字码表。

表3 由 $X = X_0 F$, (17) 式所制成的一位自动校正码

原 码 (或信息码)				校 正 码		
x_1	x_2	x_3	x_4	x_5	x_6	x_7
0	0	0	0	0	0	0
0	0	0	1	0	1	1
0	0	1	0	1	0	1
0	0	1	1	1	1	0
0	1	0	0	1	1	0
0	1	0	1	1	0	1
0	1	1	0	0	1	1
0	1	1	1	0	0	0
1	0	0	0	1	1	1
1	0	0	1	1	0	0
1	0	1	0	0	1	0
1	0	1	1	0	0	1
1	1	0	0	0	0	1
1	1	0	1	0	1	0
1	1	1	0	1	0	0
1	1	1	1	1	1	1

要解码,也很容易,最原始的解法是比较收到的信号 $\bar{X}$ 与所有 A 里的元素,取其最近者,若 $\bar{X}$ 只有一位或没有错误,则由码表

中找出这个最接近的输出码字,就是所求的校正的码字了.例如

$$\bar{X}=0101000 \tag{18}$$

把它与表1中的元素比较,可知它与0111000最近,如果只有一位错,必定是由它产生的.

但如果 m 很大, A 中有 2^m 个元素,如作一一相比较,对大量的信号而言是可观的计算量,若令 $e_0=0$, e_i =只有 i 位含1的 n 维向量,因只有一位错,则在 $\bar{X}+e_i, i=0, 1, \cdots, n$ 中必有一个属于 A .换言之,若是第 i 位错,则 $\bar{X}+e_i\in A$,因凡在 A 中之元素皆适合(17),故

$$(\bar{X}+e_i)M=0 \tag{19}$$

在计算时(19)可以简化成 $\bar{X}M=e_jM$,故只要比较 $\bar{X}M$ 与 e_jM 就可以了. e_jM 是定向量,可以先存起来,而 $\bar{X}M$ 又要乘一次就行了,若 m 不很小, $2^m\gg n$,比较 $n+1$ 个向量要比比较 2^m 向量快多了.

以表1为例, e_jM 如下表:

表4

e_j	e_jM
(0000000)	(000) (没错)
(1000000)	(111)
(0100000)	(110)
(0010000)	(101)
(0001000)	(011)
(0000100)	(100)
(0000010)	(010)
(0000001)	(001)

再以(18)为例,

$$\bar{X}M = (0101000) \begin{bmatrix} 1 & 1 & 1 \\ 1 & 1 & 0 \\ 1 & 0 & 1 \\ 0 & 1 & 1 \\ 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} = (101)$$

与 e_3M 相合，即在第3位发生错误，原码应为

$$\bar{X} + e_3 = (0111000)$$

与直接比较的答案相同。

定理5.8的推论不可以推广到 $k \geq 2$ 的情形去，因当 $k \geq 2$ 时，线性码 t 的下限与任何码的下限(11)不同。

定理5.9 若以(14)法编可以校正至多含 $2k+1$ 个数码错的码，则 t 必须满足

$$\binom{n-1}{1} + \binom{n-1}{2} + \cdots + \binom{n-1}{2k-1} < 2^t - 1 \quad (20)$$

证明 用(14)编码，我们必须找到 n 个 $t \times 1$ 的非0向量 $v_1, v_2, \dots, v_n$ 使得没有任何 $2k$ 个加起来会是0，这种向量的找法可用“筛”的方法，先取任2个异于0的向量 v_1, v_2 ，凡是与这两个向量相加可以为0的向量全部去掉，在余下的向量中任取一个作 v_3 ，筛去凡能与 v_1, v_2, v_3 加起来为0的向量， $\dots$ ，假定我们 $v_1, v_2, \dots, v_i$ 时已找到了 $2k-1$ 个向量，则我们必须对所有的 v_i ，两个 v_i 的和，3个 v_i 的和， $\dots$ 一直到 $2k-1$ 个 v_i 的和都筛去，因为它们都可以使 M 不合要求（例如我们取了 $v^* = v_1 + v_2 + v_3$ ，则 $v^* + v_1 + v_2 + v_3 = 0$ ）。这一共筛去了

$$\binom{i}{1} + \binom{i}{2} + \cdots + \binom{i}{2k-1}$$

个向量，因一共只有 $2^t - 1$ 异于0的 t 位向量，故对 $i = 1, 2, \dots, n - 1$ 而言

$$\binom{i}{1} + \binom{i}{2} + \dots + \binom{i}{2k-1} \leq 2^t - 1 \tag{21}$$

都要成立，特别是当 $i = n - 1$ 时左式最大时(21)式也得成立，故本定理得证。

满足定理5.9的 t ，经我们用计算机算了一些数据，它们的值列在表5中。可见当 $k > 1$ 时，线性码可能不是用的最小的 t 。但由于线性码容易制，而又容易解，所以是最常用的编码法。

表5 对 m 位信号码所要加的最小尾位数以作成 k 位校正码
在线性码 t 的下界由公式(20)得出，一般下界 t 由公式(11)得出

k/m	1	2	3	4	5	6	7	8	9	10	20	30	40	50
1. 线性码 t	2	3	3	3	4	4	4	4	4	4	5	6	6	6
一般 t	2	3	3	3	4	4	4	4	4	4	5	6	6	6
2. 线性码 t	4	7	8	8	9	9	10	10	10	11	13	14	15	16
一般 t	4	5	6	6	7	7	7	7	8	8	9	10	11	11
3. 线性码 t	6	11	12	13	14	15	15	16	16	17	20	22	23	24
一般 t	6	8	8	9	9	10	10	10	11	11	13	14	15	16
4. 线性码 t	8	15	16	18	19	20	20	21	22	22	26	29	31	32
一般 t	8	10	11	11	12	12	13	13	14	14	17	18	19	20

§5 结 论

有人说“好的理论，它必是简洁的”，我们已看到了校正码理论简洁的一面，其中最主要的关键在于说明距离的定义及式(17)。但我们所举的例子都在 $k = 1$ ，即改正一个错误的情形。再向前推，就应验了有些人说的：简单的东西都已经被人做

完了，所剩的全是难或繁的题目了。在 $k > 1$ 的情形，编码与解（或校正）码就要变得复杂了。一般所用的是一种叫做循环码（cyclic code）的方法，读者若有兴趣，可以在本文所附之参考资料中找到。

另外一个研究方向是针对码字中数码错成串的改正。因有时一个强烈的干扰因素可以破坏一成串的电码，这时候，本文所谈的方法就失去了效用。比如说， $m = 7$ 位的 ASCII 码加上 $t = 4$ 位的检定码可以校正一位错误，但如果遇到了一个强的干扰，可以把一串十位的数码全破坏了。当然我们再也无法用这十位数来使原信号复原了。很显然的，我们若用一个码字的传法，则我们无法恢复一个破坏得很厉害的码字。但我们若把所有码字比方说一百个码字，同时放在一起加校正码，则有可能可以恢复一长段的错误。事实上可以有一种叫Reed-Solomon的码，它只要在具833数码的信号码上加上56个检定码就可以改正不多于22个成串的数码的错误。以 ASCII 码为例，833个数码含有119个信号，而只要加56个检定码就可以改错，是非常令人惊奇的事。它比奇偶检定码还要经济，此地的主要关键在于此22个或22个以下的错误必发生在成一串的22个数码中。但它的理论也不是本文的范围可以包括的。

最后，我们要提到的是这些编码理论都有其实际用途，而且市面上有不少用这些方法制成的硬、软体装置。

小注

1. 这是假定两次错误为独立事件，有时声音会干扰一般信息，则相邻的错误就不能假定为独立事件了，在第五节内会提到这个问题。
2. 今全部码的奇偶数为1也可以。
3. 熟悉较高深数学的读者可知下面二结果

(1) $H(X, Y)$ 是一个量度(measure)。

(2) Z^n 可以想成一个交换群,或在 $\{0, 1\}$ 域中的线性向量空间。

4. 因 2^m 个码全用上了, $X \in A$, $Y \in A$, 则 $X = X_m F$, $Y = Y_m F$, 可知 $X + Y = (X_m + Y_m)F \in A$ 。

参 考 书 目

1. Djinitri Wiggert "Error-Control Coding and Applications", Artech House, Inc. Dedham, Massachusetts 1978.
2. Judith L. Gersting, "Mathematical Structure for Computer Science", W. H. Freeman and Company, San Francisco, 1985.
3. Vera Pless, "Introduction to the Theory of Errorcorrecting Codes, "Wiley & Sons, 1982.

第六章 模糊集浅介

§1 前 言

科学的进展起动于满足人类探求及解决所遭遇的问题，但它也使人们同时产生了新的欲望与好奇。数学一直是研究科学的主要工具，由开始计算量度的简单算术、三角、几何，演进到今天各式各样抽象高深理论的纯数学及各式各样的应用数学。不论怎样，所研究的对象都是在各种条件明确规定下进行的。因为人们一直根深蒂固地把数学与一丝不苟、清清楚楚的概念连在一起。这也就是说以往的数学一直建立在二值“真”与“假”的逻辑推理上。即对于一个命题的判断只用“真”或“假”来定论，不会想到渗入介中性，即亦真亦假的混合性。但大家都明白，我们在日常生活中所用的推理或言论，大多是近似推理，用一种不经细辨(模糊)的态度去求结论。过去由于人类处理资讯的容量有限，很自然使用二值逻辑，但随着电脑的普遍使用，人们觉得一个事件或概念的形成，可以把许许多多的因素都放

进来分析。比如与视觉或心理因素较有关的人文或社会学的抽象概念最好也能加以作数量化的数学讨论。尤其过去几十年来人们开始有兴趣对“人工智能学”(Artificial intelligence)谋求发展研究。人们希望对一些复杂或未完善的系统的分析与决策过程,建立一种对语法及语意有适当处理的数学模式,使在不可能很清楚(模糊的)情况下,塑造出一套可让机器“理解”一般自然语言的理论。这样把原来只能顾及到语法的奴隶性机器改进到一个可以思考的“假人”。同时也努力把原先只能顾及到语法的仙农氏(Shannon)的信息理论,改进到能处理语意的信息论。这样才能达到人对人工智能的期望及使用,也即大大提高自动化的境界。

读到人工智能其功能中一个最广的应用是图样鉴定(Pattern Recognition),简单的说是由机器代人去辨识或进行某种工作。如有些机器可以辨识人们书写的数字,大家都知人们概念中的数字决无一正正板板的形状,如“5”它可以具形状5、5、5等各式各样的情形,模模糊糊的存在人的脑海中,不同地出现在人的笔下。模糊集的概念引入使得机器在其设计及功能上能逐步地向“像人”的理想迈进。

§ 2 模糊集概念及定义

近一、二十年来自从人们把电脑广泛地应用到图样鉴定(在自动化设计中)、病情诊断、资源探测及军事上种种导向的智慧炸弹等方面后,在工程及电脑界兴起了所谓“模糊集”(Fuzzy set)的研究。到底什么是模糊集?我们其实在开场已作了伏笔,其实它也不是什么神奇出现的理论,它是一个很自然应时代而

生，随人们的产物电脑的发展引生出的一套数学理论。最早是1965年加州大学控制论教授查德（Zadeh）先生提出来的。他觉得日常的许多现象及观察结果不可能很精确的描述，比如说一个人的美丑、身材的高矮、棋艺的高明、工作的辛劳等很难定出明确的数据或标准来处理。而在人们的思维中既然有各式各样含糊性的概念在支配人们的活动，怎样将种种模糊的物类或概念用数量系或严谨的数学来表达呢？这样就促成了他对模糊集理论的创立。再想一下引入“模糊”的概念是恰当的吗？我们看看一个实际的问题，如去识别一个人，我们只要有点大概的描述，如大约多高、讲话带什么口音、戴眼镜、神情严肃、大概几岁，就不难在一个有限的范围内，完全确定这个人是谁了，不必要准确的年龄、什么牌子的眼镜、几尺几寸高。对于处理模糊概念的关键是在了解什么叫模糊集？我们知道传统的数学中所谓的一个“集”是指一个具有明确定义的事物或概念的总体，有了一个“集”我们可以用它来判断其它事物或概念是否属于该集，不会产生模棱两可的情形。

模糊集是照英文“Fuzzy set”一词翻译过来的，它其实是普通集合的一种推广。一般当我们给定一个集或集合（通常的集合意义之下），往往同时得给定一个明确的概念或具体的事物。例如集合： $A = \{1, 2, 3, 9, 17, 21\}$, $B = \{\text{世界各国的首都}\}$ 等等。因此对于任一个整数，我们可明确地判断它是否属于 A ？对任何一个都市，我们也可判定它是否属于 B 等等。换句话对于任何一事物，我们可以断然地判定是否属于某一个集。另一方面，我们可以说一个普通集必相应有一隶属的函数，即所谓的特征函数。任何集 S 中的一个子集 A ，其实就是定义在 S 上的一个特征函数 f_A （注意 f 是一般一个函数记号，把 A 记在其上是为了显明与 A 特别有关），它满足

$$f_A(u) = \begin{cases} 0 & u \notin A \\ 1 & u \in A \end{cases}$$

这表明任何 S 中一元素对子集 A 的隶属只有1(100%)与0两种极端的情形。把这种普通集合论中元素对集合的绝对隶属与绝对不隶属的关系($u \in A$ 与 $u \notin A$ 的可能),加以放松,而提出“隶属程变”的一种推广概念就是模糊集的产生。因此我们可以说模糊集是指在一已给定的讨论域 S (或定义集上,这和函数一样先要有一明确的定义域或讨论的领域)上,所规定的一个隶属程变 μ (或一函数但其值介于0与1之间,即 $0 \leq \mu(x) \leq 1, x \in S$),对任何 S 中的元素都有一个值 $\mu(x)$ 与之对应——隶属程度。所以模糊集不过是一个定义在一个普通集合 S 上的一个函数 μ ,它满足 $0 \leq \mu(x) \leq 1, \forall x \in S$ 。这样一来,所谓的普通集合是模糊集的一个特殊情形了。所以我们也可预期有些集与集间的运算结果对模糊集仍是成立的,但当然也一定有不同之处的。我们现在多少了解了对于一个模糊集,没有必要(有时也是没法子)说明它的成员(元素)是哪些,或不是哪些,而只表明一个元素能以多少的“隶属程变”(用0到1来衡量)算作为其成员。例如对天气热不热的概念是

$$\mu(m) = \begin{cases} \left[1 + \left(\frac{m-80}{10} \right)^{-2} \right]^{-1} & m \geq 80 \\ 0 & m < 80 \end{cases}$$

我们说气温 m 度(以华氏计)时算“热天”,此一模糊集的隶属程度是 $\mu(m)$,可见不到80度不算热,超过80度愈多时,就愈可被称为“热天”了。

注意:上面这个函数 μ 没有什么理论根据定出来的。

又就像在欧几里得空间中,不同单位距离定出不同的空间。在同一个定义域 S 上,不同的隶属程度(或等级函数)确定出不

同的模糊集。

还记得在前面我们用记号 f_A 表示集 S 中子集 A 所引生出的特征函数,在 S 中我们一旦定出了一个特征函数 f_A ,我们也就同时定出了一个子集 A ,反之亦然。在此意义下,我们可以认为 S 的一个子集 A 不过是定义在 S 上的一特征函数 f_A 了。对模糊集我们也作同样的认同。为示别模糊集与普通集的记号起见,我们通常在字母的下方加一个“ $\sim$ ”的记号来表示模糊集,如 $A_\sim$, $B_\sim$ 等。我们也可把 S 上一个模糊集 $A_\sim$ 与定义在 S 上的一隶属函数 $\mu_{A_\sim}$ 认同,但在不致有混淆的情形下,我们通常用记号 μ_A 来代表 $\mu_{A_\sim}$ 。

所以,我们现可叙述模糊集的定义如下:

定义6.1 在普通集合 S 上,若有一个实函数 μ_A 满足

$$0 \leq \mu_A(x) \leq 1, \quad \forall x \in S$$

则称 μ_A 决定了一个 S 的模糊子集 $A_\sim$,对所有的 x , $\mu_A(x)$ 称为对模糊集 $A_\sim$ 的隶属程变。我们也称 μ_A 为 S 的隶属函数,也就是概念 $A_\sim$ 。一个模糊集的主要关键就是如何定出与某一概念相应的隶属函数,它也是一般问题的核心。

定义6.2 若 $A_\sim$ 为集合 S 的一模糊子集,则普通集合

$$A_\alpha = \{x \mid \mu_A(x) \geq \alpha, x \in S\}$$

称为 $A_\sim$ 的 α 度普通子集。

§ 3 模糊集的运算

在实际应用问题中,我们遭遇的模糊概念可能是若干个模糊概念的组合。如聪明的年青人。所以要能对模糊集间适当地定出它们之间的运算,才能有所应用。事实上与集合的运算类似,对模糊集也可有“和: $\cup$ ”,“交: $\cap$ ”及“补: c ”等集的运

算. 这时一个很自然的想法, 对模糊集作以上的运算, 可以通过对其相应的隶属函数作布氏代数运算实现.

定义6.3 设 A, B 为集 S 上两个模糊子集 (或模糊集), 令

$$\mu_{A \cup B}(x) = \max\{\mu_A(x), \mu_B(x)\}$$

$$\mu_{A \cap B}(x) = \min\{\mu_A(x), \mu_B(x)\}$$

$$\mu_{A^c}(x) = 1 - \mu_A(x)$$

则以 $\mu_{A \cup B}, \mu_{A \cap B}, \mu_{A^c}$ 为隶属函数所相应的模糊子集分别称为 A 和 B 的“和集”, “交集”及“ A 的补集”, 并分别以 $A \cup B, A \cap B$ 及 A^c 记之, 参看下图. 并注意它们也相应于逻辑学上的“或”、“且”及“非”的三种逻辑运算.

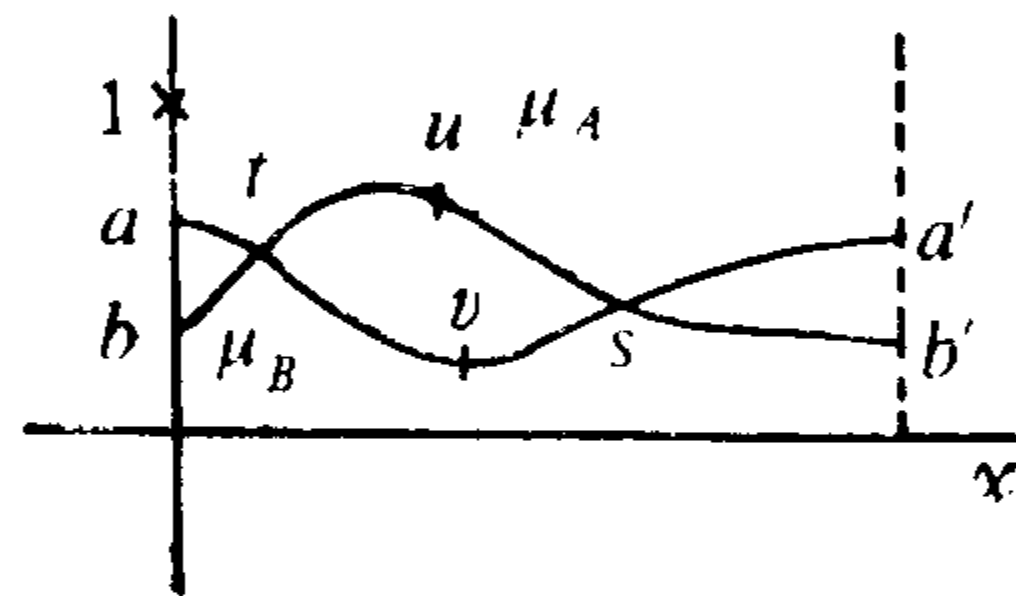


图8 由 a 经 t 经 u 经 s 到 a' 表的曲线为 $\mu_{A \cup B}$

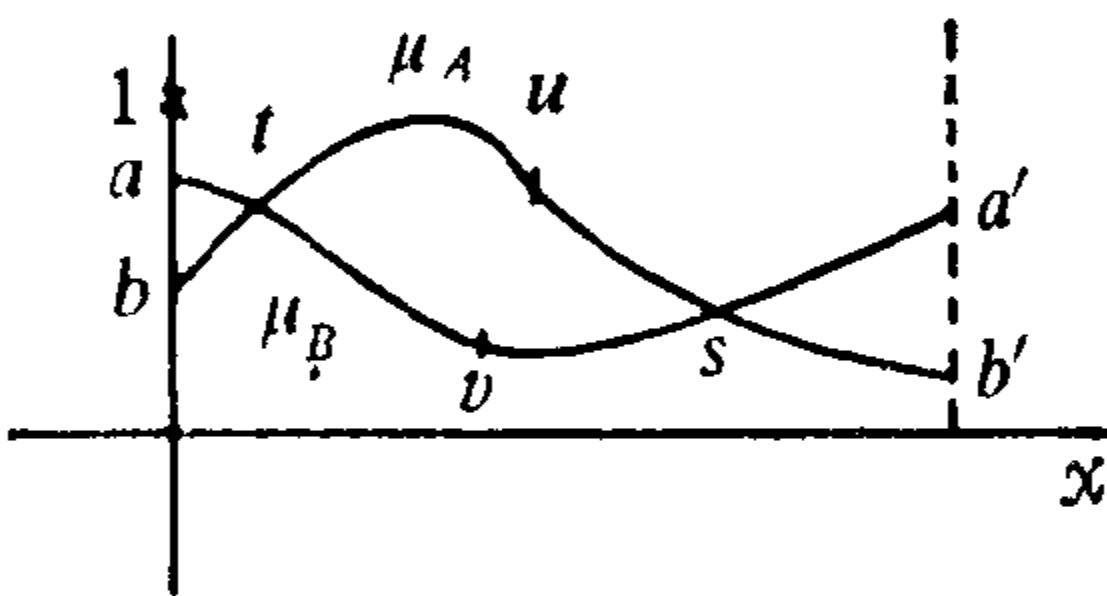


图9 由 b 经 t 经 v 经 s 到 b' 的曲线为 $\mu_{A \cap B}$

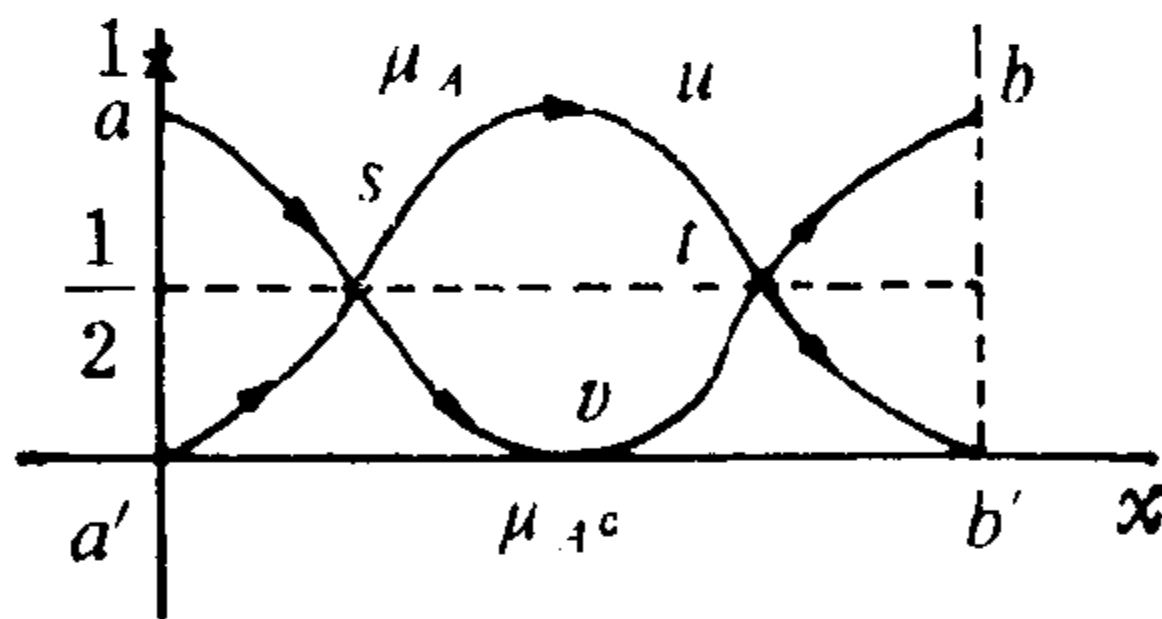


图10 则 $a' \rightarrow s \rightarrow u \rightarrow t \rightarrow b'$ 为 μ_{A^c} 的曲线
 $a \rightarrow s \rightarrow v \rightarrow t \rightarrow b$ 为 μ_A 的曲线

现在我们先看上面的定义合不合常理.

例6.1 设 S 为甲校的所有老师的集合. A 及 B 分别代表“甲

校中的年轻老师”及“甲校中的个子高老师”。两个模糊集(当然我们假定 A, B 都有了相应的隶属函数.现若甲校的某老师为“年轻”及“个子高”的隶属程度分别为0.8及0.5,那么我们自然可以认为该老师“年轻或个子高”的程度为0.8($=\max(0.8, 0.5)$)).但对同一人我们认为他“年轻又个子高”的程度只能算0.5($=\min(0.8, 0.5)$)),这是因为一般两种性质符合的可能性一定比符合单独其中之一性质的可能要低.而两种性质只要具备某一的隶属程度,当然要比具体其中一种的隶属程度来得高.

定义6.4 设 A, B 为集合 S 上的两个模糊子集,又若对所有的 $x \in S, \mu_A(x) \leq \mu_B(x)$, 则称 A 包含在 B 中, 记为 $A \subseteq B$, 若 $\forall x \in S, \mu_A(x) = \mu_B(x)$, 则称 A 与 B 相等, 以 $A = B$ 记之.

于是我们可得 $A \cup B$ 是包含 A, B 的最小模糊子集, $A \cap B$ 为含于 A 及 B 中的最大模糊子集.

由以上我们不难得到下面一系列类似普通集合运算的规则:

- | | |
|---|------|
| 1. $A \cup A = A \cap A = A$ | 幂等律 |
| 2. $A \cup B = B \cup A, A \cap B = B \cap A$ | 交换律 |
| 3. $(A \cup B) \cup C = A \cup (B \cup C)$ | 结合律 |
| $(A \cap B) \cap C = A \cap (B \cap C)$ | 结合律 |
| 4. $A \cap (B \cup C) = (A \cap B) \cup (A \cap C)$ | 分配律 |
| $A \cup (B \cap C) = (A \cup B) \cap (A \cup C)$ | 分配律 |
| 5. $A \cup (A \cap B) = A, A \cap (A \cup B) = A$ | 吸收律 |
| 6. $(A \cup B)^c = A^c \cap B^c, (A \cap B)^c = A^c \cup B^c$ | 对偶公式 |

$$7. (\underline{A}^c)^c = \underline{A}$$

我们只举对定律3(结合律)的证明, 其它可仿此.

证3 这是因对 $\underline{A}$, $\underline{B}$ 及 $\underline{C}$ 的三个隶属函数 μ_A , μ_B 及 μ_C , 对任何一 $x \in S$ 的隶属程度为 $\mu_A(x)$, $\mu_B(x)$, $\mu_C(x)$, 它们之间有下列6种情形:

$$1. \mu_A(x) \geq \mu_B(x) \geq \mu_C(x)$$

$$2. \mu_A(x) \geq \mu_C(x) \geq \mu_B(x)$$

$$3. \mu_B(x) \geq \mu_C(x) \geq \mu_A(x)$$

$$4. \mu_B(x) \geq \mu_A(x) \geq \mu_C(x)$$

$$5. \mu_C(x) \geq \mu_A(x) \geq \mu_B(x)$$

$$6. \mu_C(x) \geq \mu_B(x) \geq \mu_A(x)$$

我们只举上面六种情形任一种, 设情形1的时候来证, 其它可仿此. 这时因为

$$\mu_A(x) \geq \mu_B(x) \geq \mu_C(x)$$

$$\text{所以 } \max\{\mu_B(x), \mu_C(x)\} = \mu_B(x),$$

$$\max\{\mu_A(x), \mu_C(x)\} = \mu_A(x)$$

$$\text{及 } \max\{\mu_A(x), \mu_B(x)\} = \mu_A(x)$$

$$\text{于是 } \max\{\mu_A(x), \max\{\mu_B(x), \mu_C(x)\}\} = \mu_A(x)$$

$$= \max\{\max\{\mu_A(x), \mu_B(x)\}, \mu_C(x)\}$$

对集 S 上所有模糊子集的全体以 $F(s)$ 表之.

值得一提的是 $F(s)$ 不像普通集的子集全体在普通集合运算下形成一布尔代数, 因 $F(s)$ 不满足互补律; 若以 $\underline{I}$ 表相应

$$\mu_I(x) = 1 \quad \forall x \in S \text{ 的模糊集}$$

$\underline{Q}$ 表相应

$$\mu_O(x) = 0 \quad \forall x \in S \text{ 的模糊集}$$

则一般

$$\underline{A} \cup \underline{A}^c \neq \underline{I}$$

$$(\because \forall \mu_A, \max\{\mu_A(x), 1 - \mu_A(x)\} = 1 \quad \forall x \in S \text{ 不一定成立})$$

(3.1)

及 $A \cap A^c \neq \sim$

$$(\because \forall \mu_A, \min\{\mu_A(x), 1 - \mu_A(x)\} = 0 \quad \forall x \in S \text{ 不一定成立})$$

(3.2)

特别地，如果我们取 $A = \left\{x \mid \mu_A(x) = \frac{1}{2}\right\}$ ，则(3.1)及(3.2)就自明了。

§ 4 如何定出隶属函数

为了定出一个模糊集的隶属函数，一个似乎自然的方法是采用统计学中的抽样平均值或标准误差等理论来建造。但不能就因此认为模糊集也不过是概率或统计的一种运用罢了。因为概率上的不确定使事件的发生具有随机性，但事件本身是确定的。由于条件的不明显或不足，事件发生与否就有了多种可能性。在 $[0, 1]$ 上取值的概率分布函数就是描述此种随机性。而模糊的不确定是指元素对集合（或一个概念或定义）的隶属程度变的量度，其事件本身的含义是不确定的，而事件发生与否是确定的，在 $[0, 1]$ 上取值的隶属函数就是描述此种不确定性，我们又称之为模糊性。但要指出，有些事件的描述或表现具有随机与模糊两性质。如“明年九月在台北可能发生强大的地震”。“强大”是对地震的一个不确定的描述，为一模糊的概念，而加上明年九月发生这是对出现地震事件的不确定的说法具有随机性。总之，统计是把总结必然的现象扩大到与偶然现象的研究，而模糊集是把清晰的现象讨论推展到迷糊的现象中去。

至今为止，除了特殊情形外，一般我们还没有一个理论性的方法来建造一个模糊集的隶属函数。对不同的问题及概念，我们往往凭一些直觉的体验或经验定出一个适当的数学隶属函数式子来供使用。下面我们举一些例子，说明对一些模糊集如何适当地给出它们具体的隶属函数来。

例6.2 设 S 为所有的三角形的集合，并且我们对相似三角形不加区分。所以这时我们依据三角形的三个内角来定义出一些模糊集如下：

首先，我们可以先按三角形的三个内角的变数以 α, β, γ 记之，并设 $\alpha \geq \beta \geq \gamma$ 。因此 S 中每一个元素（即一个三角形）相应于一序对 $u = (\alpha, \beta, \gamma)$ ，所以 S 又可表为： $\{u = (\alpha, \beta, \gamma) | \alpha + \beta + \gamma = 180, \alpha \geq \beta \geq \gamma \geq 0\}$ 。

我们现对 S ，考虑三个模糊集：

(1) “近乎正三角形”的模糊集 E 。

这时我们似乎可从许多不同图形观察，得到一个共同的印象，即最大角与最小角的差愈小，则愈像正三角形。因此我们不妨定：

$$\mu_E(u) = 1 - \frac{\alpha - \gamma}{180}$$

此一隶属函数反映了当 u 为一正三角形， $\mu_E(u) = 1$ ，即不会判断错误。

(2) “近乎直角三角形”的模糊集 R 。

这时一个如此三角形的特征是最大角度与90之差愈小，则愈像直角三角形。因此我们不妨定义隶属函数如下：

$$\mu_R(u) = 1 - \frac{|\alpha - 90|}{90}; \quad \mu = (\alpha, \beta, \gamma) \in S$$

(3) “近乎等腰的三角形”的模糊集 $\underline{G}$ 。

一个三角形的三内角如有两个角的差愈小，则愈像等腰三角形。这时我们不妨定隶属函数如下：

$$\mu_G(u) = 1 - \frac{\min(\alpha - \beta, \beta - \gamma)}{60}, \quad u = (\alpha, \beta, \gamma) \in S$$

我们也可借上面的模糊集看看一些它们间的逻辑运算结果。设 $\underline{E}$ ， $\underline{R}$ 及 $\underline{G}$ 如上面所述的三个模糊集，则 $\underline{A} = \underline{R} \cap \underline{G}$ 表达“近乎等腰直角三角形”的模糊集，及

$\underline{B} = (\underline{R} \cup \underline{G} \cup \underline{E})^c$ 表达“无什么特殊性质的三角形”的模糊集。

§5 模糊集的一些应用

5.1 模糊分类原则(模型识别的一方法)

一般设一个集 S 上定有 n 个模糊集 $\underline{A}_1, \underline{A}_2, \dots, \underline{A}_n$ 。现如果 u_0 为 S 中的一元素，我们要判定 u_0 到底归属于哪一个模糊集呢？换句话说，我们如果想知道 u_0 与 $\underline{A}_1, \underline{A}_2, \dots, \underline{A}_n$ 中哪一个最近？一个自然而直觉的判断法是比较： $\underline{A}_1(u_0), \underline{A}_2(u_0), \underline{A}_2(u_0), \dots, \underline{A}_n(u_0)$ 中的值，将 u_0 归具最大值的集，即若 $\underline{A}_1(u_0) = \max(\underline{A}_1(u_0), \underline{A}_2(u_0), \dots, \underline{A}_n(u_0))$ ，则判定 u_0 隶属于 $\underline{A}_1$ 。

例如在前面举的三角形集合 S 中，共有5个模糊集， $\underline{E}$ ， $\underline{R}$ ， $\underline{G}$ ， $\underline{A}$ 及 $\underline{B}$ 。现如 S 中的一元素 $u_0 = (85, 50, 45)$ ，则由计算可得 $\underline{E}(u_0) = 0.7$ ， $\underline{R}(u_0) = 0.916$ ， $\underline{G}(u_0) = 0.94$ ， $\underline{A} = 0.916$ 及 $\underline{B} = 0.06$ ，其中 $\underline{G}(u_0) = 0.94$ 为最大，故按分类原则，宜判定

u_0 为近乎等腰三角形。

5.2 符合变、择近原则及聚类分析

上面提的分类是以针对集合 S 中个别元素来作判别，但在普通一般较实用及复杂的识别问题中，所要识别的对象是 S 中一个子集（普通的或模糊的）。比如对某校所有的学生(S)书法体“像颜体”的形成一模糊子集 $\tilde{Y}$ ，“像柳体”的形成一模糊子集 $\tilde{L}$ ，“像于体”的也形成一模糊子集 $\tilde{X}$ ，现如何比较此三个模糊集的相近性呢？这时所要关注的是两个模糊子集间的符合程度。先介绍一个定义如下：

定义6.5 设 $\tilde{A}, \tilde{B}$ 为集 S 上的两个模糊子集，则

$$\tilde{A} \circ \tilde{B} = \max_{u \in S} \{ \min(\mu_{\tilde{A}}(u), \mu_{\tilde{B}}(u)) \}$$

$$\tilde{A} \oplus \tilde{B} = \min_{u \in S} \{ \max(\mu_{\tilde{A}}(u), \mu_{\tilde{B}}(u)) \}$$

分别称作 $\tilde{A}$ 与 $\tilde{B}$ 的内积及外积。值得提的是此两运算规则 $\circ$ 及 $\oplus$ 相当于 $\cup$ 及 $\cap$ ，所以不但满足第3节中的各运算律，且满足互补律。

定义6.6（符合度）定

$$(\tilde{A}, \tilde{B}) = \frac{1}{2} [\tilde{A} \circ \tilde{B} + 1 - (\tilde{A} \oplus \tilde{B})]$$

为 $\tilde{A}$ 与 $\tilde{B}$ 的符合程度。

在农作物品种改良时，如何判定新的品种与旧有的相近比较、医学上癌细胞的识别及其它各种研究实验上都可用到这一原则。

择近原则

设 S 上给出了 n 个模糊子集 $\tilde{A}_1, \tilde{A}_2, \dots, \tilde{A}_n$ 及另一模糊子集 $\tilde{B}$ ，若有一 $i, 1 \leq i \leq n$ 使得

$$(\underset{\sim}{B}, \underset{\sim}{A_i}) = \max_{1 \leq i \leq n} (\underset{\sim}{B}, \underset{\sim}{A_i})$$

则称 $\underset{\sim}{B}$ (在 $\underset{\sim}{A_1}, \underset{\sim}{A_2}, \dots, \underset{\sim}{A_n}$ 之中) 与 $\underset{\sim}{A_i}$ 最符合。

本节最后要提的是模糊集在聚类分析上的作用。

从数学上来讲，一个明确的分类，是要由一个等价关系来确定的(集合 S 的一关系 R 如果对任何 $x, y, z \in S$ ，满足(1) xRx (反身性)，(2) 另要 xRy 就有 yRx (对称性)，(3) 只要 xRy 及 yRz 就有 xRz (传递性)，就称 R 为一等价关系。例如在实数上“=”的关系就是一等价关系，但“ $\geq$ ”就不是了)，因此 x, y 等价就表示此两元素在某种意义下“相等”，比如说三角形的“相似”关系，学生间的同班关系。有了等价关系就可把原来的集较明显的分成许多不相交的子集。不同子集就相应不同的类，这就叫聚类分析。但现实问题的分类多半随着模糊性，因此如果把大致相似或相近的元素合成一类的聚类概念，也自然地要靠一种“模糊关系”概念来处理了，下面我们就讨论此种关系了。

模糊关系

大家也许一定记得当 U, V 为两个普通集合时，两者的笛卡儿乘积：

$$U \times V = \{(u, v); u \in U, v \in V\}$$

它代表的是任何一个 U 中的元素与任何一个 V 中的元素间所有可能的配合。如果我们对这种配合加了某一种限制，那么这种限制就造成了 U 与 V 间的某种特殊关系。反过来说，任何一种特殊关系 (或规则) 必相应 $U \times V$ 的一个子集 (即我们可以得到所有的序对 (u, v) ，其中 u, v 满足所定的规则或关系者)。对于模糊集，我们也就可以有类似的定义了。

定义6.7 设 U 与 V 为两个集合 (U, V 不一定为不同的集

合)。所谓 U 到 V 的一个模糊 R 关系是指 $U \times V$ 的一个模糊子集 (注意因 $U \times V$ 亦为一个集合, 所以可以讨论它的模糊子集)。隶属程度 $\mu_R(u, v)$ 表示 u 与 v 具有关系 R 的程度。

在 U 与 V 皆为有限集时, 我们往往可用矩阵来表示隶属程度。

例6.3 今设 S 为一组7个小孩的集合, $U = (3.5, 3.6, 3.7, 3.8, 3.9, 4.0, 4.1)$, $V = (60, 70, 75, 80, 85, 90, 95)$ 分别表示 S 中小孩的身高 (以呎计), 及 S 中小孩的体重 (以磅计)。则 U 与 V 之间的一个模糊子集 R 的隶属程度可以表示如下: (即 R 表示了身高 $\longleftrightarrow$ 体重间的一模糊对应)

$\mu_R(u, v)$ $U \backslash V$	V						
	65	70	75	80	85	90	95
3.5	1	0.7	0.4	0.2	0	0	0
3.6	0.9	1	0.8	0.7	0.4	0.2	0
3.7	0.7	0.8	1	0.8	0.7	0.6	0.2
3.8	0.5	0.7	0.8	1	0.8	0.5	0.3
3.9	0.4	0.5	0.7	0.8	1	0.8	0.6
4.0	0.2	0.3	0.6	0.8	0.9	1	0.4
4.1	0	0.2	0.5	0.7	0.8	0.9	1

此相应的矩阵又叫模糊矩阵。

有了模糊关系, 我们也可利用它们的模糊矩阵定义它们间的一些运算 (也即把前面内、外积的定义推展到高维空间而已)。有了模糊关系, 我们就可定义模糊等价关系, 并进而可定出信靠度及影像度等量度, 配合模糊算法语言, 就可在利用电脑整理大批资料或文件时, 能动态地将所要的资料一层一层地自动

滤出来。

结论：模糊集的概念已是现代管理、控制自动化利用电脑操作中仿人脑功用的数学理论根据。目前它的理论研究配合有概率论、逻辑、抽象的拓扑学、代数学及测度与空间理论的各种架式，我们在此介绍的只是模糊论最浅显的一面。坊间已有不少此类的专书(用英文写的)，但在台湾似还未看到中文写的有关的书，希望不久就可看到它的出现。

附注

现在日本工商业界已将模糊理论应用在家庭电器用品的发展及制造上，取得了相当惊人的进展，如日本已出售一种可随室内光线的明暗，自动调节其画面明暗的“模糊电视机”，该机且能随观看电视的人与画面的远近有所改变，会在其音量上作出适当的调节，“模糊洗碗机”可视其收进的碗碟的脏度及多少，来决定其冲洗的时间、洗涤剂、水量。另外类似的也有所谓的“模糊空调机”。这一切都说明将来人们对其器具的发明及环境的适应，都要以较精确(多层面)的模糊理论来控制了。

参 考 文 献

1. 杨重骏、杨照崑合著“现代应用数学”，台北，东华书局，1984。

第七章 机率程式与随机数

§1 前言

一般计算机程式都是一步接着一步，什么样的输入，一定会产生什么样的作业运转及什么样的输出，不论在哪里跑这个程式的进行及结果都该是一样。

机率程式 (probabilistic algorithm) 的意思是指程式中含有不确定的元素，造成运转并非一成不变，同样的输入，可以产生不同的运算，但结果仍然是所要求的结果。基本就来，这种程式设计往往用到统计学的原理，在某种情形下，我们若对输入资料的结构有相当的了解，就可以大大的节省计算的时间，先举一个简单的例子。

假定我们有十万个缴税单在磁碟中，我们要取出前面一千个最高额的税户，一个直觉的方法就是先取最高税户，再取次高税户，…，一直取到第一千个高税户为止，如何找到最高税户？我们必需将所有的税额都互相比过（少一户都不行），我

们若每次比较都取大户再往下比,要比9999次,一次也少不了,若再往下去,我们得比9998次去找到第二大户,……。若令 $N = 100000$,则找到 $n = 1000$ 最高税户的总比较次数是

$$S = (N - 1) + (N - 2) + \cdots + (N - n) = nN - n(n + 1)/2$$

因 N 比 n 大得很多,故 $S \approx nN = 10^8$ 。

但是如果我们知道至少有1000户的税额大于 $M = 200000$ 元,则我们可以先把所有的税与 M 比较,凡是小于 M 的税额就不必再用,这时我们用了 N 个比较,假定我们发现 $N' = 2000$ 个税户大于 M ,那么我们现在再在这 N' 户头中找一个最大户,若用刚才同样的方法,则我们需要

$$S' = nN' - n(n + 1)/2 \approx 1.5 \times 10^6$$

个比较,再加上原先的 $N = 10^5$ 个比较,全部的比较是 1.6×10^6 ,是原先 $S = 10^8$ 的六十二点五分之一,相当于一分钟与一小时计算之差。

不过在实际作业中,我们可能不知道 M ,因此我们必须要有个方法来猜测 M 之值,我们所希望的 M 是指一千个大户的下限,我们可以定它为从最高额向下数第一千个大户的税额,目前测定 M 最好的方法是有统计抽样的方法,即在原数据中抽出若干样本,用样本的分布情形来估计 M 。直觉上,我们可用样本的前百分之一(n/N)的下限 M' 作为 M ,但我们又可以觉察到这个 M' 大于 M 与小于 M 的机会像是一半一半,因此有一半的可能大于 M' 的户头小于一千,那么先取所有大于 M' 的户头的方法就不能产生一千个大户,因此我们不妨保守一点,用样本的百分之二的下限 $\hat{M}$ 为 M 之估计值。现在我们要问在原数据中大于 $\hat{M}$ 者约有多少?要求出这个值,我们必须知道原来的抽样是随机抽样,即每个税户都有相同的机会(n/N)被抽到。

设原数据中没有重复,因此我们可以令所有的数据为

$$x_{(1)} < x_{(2)} < x_{(3)} < \cdots < x_{(N)}, \quad (1)$$

则大于 $\hat{M}$ 的税户为 i 户，即 $\hat{M} = X_{(i)}$ ，其或然率为

$$P_r\{\hat{M} = X_{(i)}\} = \frac{\binom{i-1}{y-1} \binom{N-i}{n-y}}{\binom{N}{n}}$$

式中 $n=1000$ ，即抽样之量， $y=20$ 即样本之前百分之二十。

若是 $\hat{M} \geq X_{(n)}$ ，我们的目的就达到了（因为若以 $\hat{M}$ 为下限，所有大于 $\hat{M}$ 的数据大于 $n=1000$ ），其或然率为

$$p = P_r\{\hat{M} \geq X_{(n)}\} = \sum_{i=n}^N \frac{\binom{i-1}{y-1} \binom{N-i}{n-y}}{\binom{N}{n}}$$

这个或然率值不太容易计算，但若用比较高深数学上的近似公式可得 $p \approx 0.9992$ ，也就是说如果用 $\hat{M}$ 为下限而取所有大于 $\hat{M}$ 的数据，但其中含不到一千个最大的税户的机率是微乎其微（ $=1-p=0.0008$ ）。

当然有可能 $\hat{M}$ 比 M 小得很多，在这种情形下的 M 就不算对 M 很有效的估计，但无论如何 $\hat{M}$ 不太可能很小，因此应该比不知道 M 好，在此我们不去讨论 $\hat{M}$ 与 M 之间关系的细节，只想说明我们有可能测定 M 及用测定的 M 值来节省计算的时间，因每次抽样的结果并不相同，程式运转的情形也不尽相同，但所达到的目的是相同的。

§ 2 快速编排

编排是指把数据依某一指定的次序排列，例如电话号码簿按姓名笔划排，联考成绩按总分多少排。编排的程序在计算机

程序中占重要的地位，因为它与寻找一个指定的数据有密切的关系。试想若电话号码簿不按姓名笔划排，则要找一个人的电话是何等的困难，又如上节的例子，如果税单是照税额大小排的，则我们很容易找到第十号大税户，否则就要排很多次才能找到；由于数据不断的更新，计算机内的资料需要不断的编排，因此几乎所有计算机程式中都有一个编排程式，而其中最常用的叫做快速编排(quick sorting)，是 Hoare 在1962年发现的。

快速编排的原理与上节所讨论的抽样法相似，但比较简单。

现假定我们有 n 个不相同的数据，我们要把它们按大小排起来，刚才说的直觉方法是先比较 $(n-1)$ ，次找最大的，再比较 $(n-2)$ ，次找次大的， $\cdots$ ，一直到比较了

$$\begin{aligned} T &= (n-1) + (n-2) + \cdots + 3 + 2 + 1 \\ &= n(n-1)/2 \approx n^2/2 \end{aligned} \quad (2)$$

次之后，就编排完成。

现在假定我们知道这组数据的中间值（即一半比此数大，一半比此数小），则我们可以先将每个数据与此数比较，若小于它，放在一堆，大于它的放在另一堆，这样一共用了约 n 个比较，然后再用前面的方法把大小两堆分别按大小编排再合起来，很显然的，我们用了

$$n + 2((n/2)^2/2) = n + n^2/4 \quad (3)$$

个比较。上式中 $(\frac{n}{2})^2/2$ 是大小两半编排时所需用的比较数，

因有二半所以乘以 2，比较(2)式与(3)式，可知在 n 很大时， $n + n^2/4 \approx n^2/4$ ，省了几乎一半的时间，以此类推，我们若又能知道上半或下半的中间值，我们又可以省一半的时间。

现在问题是如何找到这个中间值，根据上节的想法，我们可以取一个样本，用以测定中间值。取多少样本？Hoare 的建

议是取一个样本，直觉上看来这个样本实在太小。但出乎意料之外的，这还是目前最好的办法，有人曾经想加大样本以对中间值作更精确的估计，但结果因程式变得更复杂，并不能节省什么计算时间。因此大家仍用 Hoare 的方法随便取一个值做中间值。

因机率程式的运转与抽出的中间值有关运转的时间就不是一个定数，一般我们要考虑二种情形：

一、运转时间的平均值(期望值)。

二、在最坏的情形下，运转的时间。

其中第二种情形是要防止在最坏的情形下，程式不会永远跑不完。

在编排程式中，运转时间几乎与比较次数成正比，因此我们就以比较的次数来计算运转时间，先看期望值。在快速编排中，令

$S(n)$ = 一组含 n 个不同数据用 Hoare 方法所需之平均比较次数

用(1)式之符号，我们有相同的机会 $(1/n)$ 取 $X_{(i)}$ ， $i = 1, 2, \dots, n$ 为中间值，设我们取出之值为 $X_{(m+1)}$ ，则原数据被分为 m 与 $n - m - 1$ 两部分，故

$$S(n) = n - 1 + \frac{1}{n} \sum_{m=0}^{n-1} \{S(m) + S(n - m - 1)\} \quad (4)$$

上式中 $n - 1$ 是指剩下的 $n - 1$ 个数与中间值(估计的) $X_{(m+1)}$ 相比较以分开成小于 $X_{(m+1)}$ 的集团(含 m 个数据)及大于 $X_{(m+1)}$ 的集团(含 $n - m - 1$ 个数据)的比较次数。

式(4)很容易写成

$$S(n) = n - 1 + \frac{2}{n} \sum_{m=0}^{n-1} S(m) \quad (5)$$

$$= n-1 + \frac{2}{n}S(n-1) + \frac{2}{n} \sum_{m=0}^{n-2} S(m) \quad (6)$$

将上式最后一项代入(5)式的另一个形式

$$\sum_{m=0}^{n-1} S(m) = n[S(n) - (n-1)]/2$$

可得

$$\frac{S(n)}{n+1} = \frac{S(n-1)}{n} + \frac{4}{n+1} - \frac{2}{n} \quad (7)$$

将(7)式一直代下去可得

$$\begin{aligned} \frac{S(n)}{n+1} &= \frac{S(n-2)}{n-1} + 4\left(\frac{1}{n+1} + \frac{1}{n}\right) - 2\left(\frac{1}{n} + \frac{1}{n-1}\right) \\ &= \dots \\ &= \frac{S(1)}{2} + 4\left(\frac{1}{3} + \frac{1}{4} + \dots + \frac{1}{n+1}\right) \\ &\quad - 2\left(\frac{1}{2} + \frac{1}{3} + \dots + \frac{1}{n}\right) \end{aligned}$$

$S(1)$ 显然为0, 若以 $\ln(n)$ 表 n 之自然对数, 则一个有名的公式

$$1 + \frac{1}{2} + \frac{1}{3} + \dots + \frac{1}{n} \approx 0.577 + \ln(n)$$

可使得

$$\begin{aligned} \frac{S(n)}{n+1} &\approx 4(\ln(n+1) + 0.577 - 3/2) \\ &\quad - 2(\ln(n) + 0.577 - 1) \end{aligned}$$

即

$$S(n) \approx 2n \ln(n) \quad (8)$$

所以我们的结论是快速编排平均要用 $2n \ln(n)$ 个比较, 在 n 很大的时候 $2n \ln(n)$ 与 $n^2/2$ 相差极大, 比如说 $n=10^8$, 即一百万个数据, 则

$$2n\ln(n)=27.6\times 10^6, \quad n^2/2=5\times 10^{11}$$

若一个计算机每秒可以比较 10^6 次，则上两项的时差是 27.6 秒与五天十九小时。

快速编排最坏的情形是每次都选到最大的或最小的做中间值，在这种情形下，等于没有分，读者可以很容易求出其所需之比较次数与(2)式相同，但其可能性为

$$\frac{2^n}{n!} \approx \left(\frac{2e}{n}\right)^n \tag{9}$$

式中 e 为自然对数的底，只要 n 稍微大一点，(9)式所表示的或然率可以说是不可能。

我们曾用小电子计算机作了一个实验，在实验时我们写了两个程式，一是先取最大，再取次大的直觉方法，称之为逐次编排，另外是快速编排，在表1中有当 $n=100$ 及500的编排时间，及与理论值 $(n^2/2)$ 及 $(2n\ln(n))$ 的比较。

表1 两种编排的比较

n	快速编排 (秒)	逐次编排 (秒)	快排/ $(2n\ln(n))$	逐排/ $(n^2/2)$
100	6	25	0.0065	0.005
500	38	620	0.0061	0.005

可见我们的预计值 $(n^2/2)$ 及 $(2n\ln(n))$ 相当的正确，其中比值有0.0065及0.005之差，表明快排的程式较逐排为复杂，所用的非比较运算较多。

快速编排是否可以改进？答案是还可以，但进步不了多少，原因是这样一个直觉的定理(并不容易证明)，即比较一次只可以辨别一个元素属于二个子集中的一个，举一个例子，如果有一数 x ，我们若把它与50比较，我们只知道它是大于50或小于等于50，不可能确知 x 是何数，但如果一个集合中只有两个元素，

一个大于50，一个小于50，则只要比较一次就可知道 x 是哪一个元素，但若某集合 A 中有三个元素，则不能比一次就可确知 A 中的某一元素 x ，必须比较两次，依此类推可以知道。

定理 若 A 中含有 k 个元素，则至少要比 $\log_2 k$ 次才可以确知 $x \in A$ 是哪一个元素。

编排 n 个数据，相当于知道原来 n 个数据的排列，因 n 个元素有 $n!$ 种排列法，故至少要用

$$S_{\min} = \log_2(n!)$$

次比较，用Stirling公式 $n! \approx \sqrt{2\pi n} n^n e^{-n}$ 可知

$$\begin{aligned} S_{\min} &= \log_2(n!) \approx n \log_2 n - n \log_2 e \approx n \log_2 n \\ &= 1.38n \ln(n) \end{aligned} \tag{10}$$

将(10)与(8)比较可知 $S_{\min}$ 与 $S(n)$ 的差别已不算大，因此除非有一个相当简单的方法，就不值得用来代替Hoare 的快速编排。

§3 随 机 数

在上节谈到的Hoare方法中我们需要在 $(1, 2, \dots, n)$ 个元素中任意取一个，如何任意取？照定义，即是任何一个数都有 $1/n$ 的机会被取到，设我们可以找到一个在 $(0, 1)$ 之间的随机数 U ，它在 $(0, 1)$ 区间均匀分布着，即

$$P_r\{a < U < b\} = b - a, \quad 0 < a < b < 1 \tag{11}$$

那么我们取一个 U ，若 $0 < U < \frac{1}{n}$ 取1， $\frac{1}{n} \leq U < \frac{2}{n}$ 取2， $\dots$ ， $\frac{i}{n} \leq$

$U < \frac{(i+1)}{n}$ 取 i ， $\dots$ ， $\frac{(n-1)}{n} \leq U < 1$ 取 n ，则每个数字 $1, \dots, n$

都有 $\frac{1}{n}$ 的机会被取出。一个较简单的分配法，是令 $[x]$ 表示不小

于 x 的最大正整数，即 $[1.5]=1$ ， $[2.6]=2$ ， $[3]=3$ ，则在1，2， $\dots$ ， n 所取的随机整数可用

$$J=[nU+1]$$

随机数除了(11)式的要求外，尚需要所产生一连串的随机数之间为互相独立，若我们产生了两个随机数 U_1 、 U_2 则必须有

$$\begin{aligned} &P_r\{a_1 < U_1 < b_1 \text{ 且 } a_2 < U_2 < b_2\} \\ &= P_r\{a_1 < U_1 < b_1\} P_r\{a_2 < U_2 < b_2\} \end{aligned}$$

换言之，

你给了我 U_1 ，我不可能用它来预测 U_2 。

用计算机来产生随机数必须有一个程式。但有了一个固定的程式，随机数都定了。当然在理论上已不成为随机了。但因为对不知该程式的人而言，看了这个数预测不到下个数就算可以了。当然要严格地讲起来，这些只能算是人造随机数。

产生随机数的方法之一像轮盘赌中轮盘摇出的号码，当你把轮盘一转，无法预测指针会停在什么位置，因为用力只要差一点，停止的位置就会差得很多，轮盘转一圈又回头了，这就相当于整数论中的同余观念，如果一个轮盘分100格，那么101格就与第1格相重合了，因此在同余的观念下，对100而言

$$121 \equiv 21, 1156 \equiv 56$$

我们用(mod 100)作记号，即

$$X \equiv Y \pmod{m}$$

表示 X 为 Y 除以 m 的余数， $0 \leq X < m$ ，一般产生随机数的方法足取三个正整数 a 、 c 与 m ，用

$$I_n \equiv aI_{n-1} + C \pmod{m}, n=1, 2, \dots$$

重复产生 I_n 。当然我们需要一个开始的种子数 I_0 。(这可任意取定的)，而随机数即定为

$$U_n = I_n/m$$

在此要注意的是 a 、 m 、 c 不可以随便取，要经过多次实验合乎随机的要求，像IBM/370所用的是 $a=7^5=16807$ ， $C=0$ ， $m=2^{31}-1$ ，一般小计算机现在都有自己的随机数程式，甚至小的科学计算器也有随机数的程式，若要自己产生随机数，下面三个数可以用

$$a=25173, c=13849, m=65536$$

随机数的另一用途是做含有或然事件的模拟实验，举一个例子，现在世界少、青少及青棒球赛都采取双淘汰制，设现在四队，各队互相之间的相对实力如表2所示，表中数字表示上队(阿拉伯数字)取胜直队(中文数字)的或然率，举例而言，二

表2

	2	3	4
一	0.6	0.4	0.5
二		0.7	0.4
三			0.3

队胜一队的或然率为0.6，三队胜二队的或然率为0.7，若要根据表2去计算各队赢得冠军的或然率非常困难，但我们可以作一个模拟实验，即每次比赛时取一个随机数，来决定两队的输赢，比如二队同一队打球，我们先订取一随机数，若此数在0到0.4之间算一队赢，否则二队赢，这就符合了二队胜一队或然率为0.6的要求，现在我们就取一随机数就可以定两队的输赢，用计算机，全部的球赛一下子就打完了，我们就知道谁是冠军，我们可以打一千次或一万次，看看各队得冠军的数目，也就推测出了各队得冠军的或然率。

现代有许多决策都与不可预测的事件有关，像台风、地震、

火灾、车祸等等，但我们往往知道这些事件发生的或然率，那么我们就可以做模拟实验，看看我们决策成功的可能性，带有随机数的模拟实验已是现在许多行业中不可少的工具了。

§4 结 论

计算机日益普遍，因此也产生了很多有关计算机的笑话及漫画，也有人把一切责任推到计算机身上，我们常听到有人说：“这不是我的错，计算机打出这样的结果，我有什么办法？”，其实计算机只能做人所教它做的，一点也不多，一点也不少，这是计算机的优点，但也是计算机的缺点，计算机没有创造性。有人说，如果有人能教计算机写出一些笑话，‘人工智慧’（或人工智能）就成功了，这话十分有道理，笑话都是一种预料不到的结果或结论，若一切都在预料之中，我们就不会有笑话的感觉了。

我们相信用计算机写笑话一定与机率程式有关，即程式运转中含有写程式者不能预料的结果，使你不得不大呼这个程式真是高明，我们也相信在人类的思维中，也有一种不可确定的路径，东冲西撞，一条新的定理或笑话就出现了，但这种不可预定的尝试似乎又不是乱冲乱撞，像是在不可预料的乱撞中有一个次序，到底是什么回事，可以说是人类未来最有趣的挑战。

第八章 电传签字

§ 1 前 言

到现在为止一般人在合约或公文，或支票上的签字及图章来作为当事人的签署保证。但这往往要靠对方用眼睛及查对档案来肯定其真实性，这不是很可靠的。何况图章及签字都可伪造仿真的（事实上就是同个人的签字也每次都稍有不同！）。所以照目前一般防范伪造及冒名签字的技术来讲，要想确定一个经过电传的信息是否由信中所自称发信的人所拍发更是难上加难的现实问题了。当然，如时间及环境许可，一些重大的交易或极机密的指示可以借助电话的联系作沟通，或再拍送电讯请对方当场再作保证，或由密语来肯定。但这种往返的鉴定，不符合现时代商业上的交易（或军事上的行动指示），其讲求快速及保密的原则。所幸现有了我们在先前的一些章节介绍利用电脑来计算的公开密码的原理及其长处。发现它也可用来作电传签字的鉴定，此方法当然也有被冒签的可能，但其可能性非

常的渺少，不幸的是这种鉴定法也有缺点，即它尚不能普及地施用，只有少数大银行或情报机构有能力来进行这种电传签字的技术及鉴定，但其理论仍是很值得玩味的。

§ 2 公开密码原理的回顾

公开密码有若干不同的形式(请参看第二章)。其中之一是利用一个大的合成数 m 来作为发码的模数，通常 m 是取为两个很大质数 p 、 q 的乘积，即

$$m = p \times q \quad (1)$$

目前的实用上通常取 p 、 q 至少为50位数，因此 m 也就近至少为一个100位数了。前面我们也曾指出在目前电脑计算速度及能力的限制下，是不可能在几年内做到将 m 的因子找出来的(除非 m 具有特别的形式！)。

对 m 我们相应有一个尤拉数 $\phi(m)$ 是指所有小于 m 而与 m 互质(包括1)的正整数的数目。特别当 $m = pq$ ， p 、 q 为质数时， $\phi(m) = (p-1)(q-1)$ ，现公开码需要一个发码指数 a 及一个解码指数 d ，通常 a 是取一个相当大的数，但与 $\phi(m)$ 互质，且不大于 p 及 q ，即

$$(a, \phi(m)) = 1, a < p, a < q \quad (2)$$

以下以 ω 表 $\phi(m)$ 。

而 d 是相应 a 满足

$$da \equiv 1 \pmod{\omega} \quad (3)$$

的一正数。由上式子可得

$$d \equiv a^{\omega-1} \pmod{\omega} \quad (4)$$

这个数 d 是要保密的，也假设只有收方知道的一个数！

$$(m, a) \quad (5)$$

称为拍码组，这是公开的。所谓公开密码是拍送电码的一方甲在拍发电码给要收的一方乙时，不必对其所要发出的电码保密。甲只要照收方乙预先发布的拍码组，把要拍的数码 M 以

$$E \equiv M^a \pmod{m} \quad (6)$$

公开地拍送到收方乙。当乙收到后就只有他本人知道的解码指数 d ，对正进行解码(或译码)。其方法如下：

计算 E^d

$$E^d \equiv (M^a)^d \equiv M^{ad} \equiv M^{a d + 1} \pmod{m} \quad (7)$$

($\because ad \equiv 1 \pmod{\omega} \implies ad = 1 + k\omega$, k 为一整数)

但由Euler定理

$$M^{\omega} \equiv 1 \pmod{m}$$

故(7)式可导至

$$E^d \equiv M \pmod{m}$$

于是乙收到甲拍出的数码 M 了。

所以除非有乙方的解码指数 d ，要将 E 解回得 M 几乎是不可能的，而要想试得 d 根据(4)，必须知道 ω ，即 $\omega = (p-1)(q-1)$ ，也即 p 及 q ，现虽然我们知 $m (=pq)$ ，但要在短时间内将 m 分解出 p 、 q 两因子是几乎不可能的，这也是以上所用的公开密码的关键。

§ 3 电传签字系统

我们这儿的电传签字，就是甲在信息中(提及自己名字)在信息尾端附上代表其姓名(或加上住址、或机关、行号)的数码(这是公开的)，但甲利用其解码指数将此数码加以重新编码(签

字), 然后照收方乙公开公布的发码组, 把信息(包括签字)重新编拍给对方乙。在乙收到信息后, 先用其自己的解码指数解出信息。如乙发现其中提了甲的名字就知道信尾的数码为姓名的签字了。这时乙方就用甲公开的发码组的发码指数对电传的签字(数码)作鉴定, 看是否和信中所提的甲的姓名相符。这就是电传签字的鉴定了。

现在我们看电传签字的一例子及数学理论过程。

设有一个人甲在外远离其银行乙所在地, 急需乙汇寄一大笔钱。设甲的公开发码为 (a_0, m_0) , d_0 为其解码指数, (a_1, m_1, d_1) 为乙的码组。于是甲用电传写了个信息(其中提及他的名字 N)并在信息后把他姓名数码 N 作了签字, 即数码

$$T \equiv N^{d_1} \pmod{m_0} \quad (8)$$

将此附在信息数码 M (其中提及甲的姓名 N)之后, 照乙方的发码组(这是公开的)作

$$[M \cup T]^{a_1} = [M]^{a_1} \cup T^{a_1}$$

拍送到乙方。在乙收到后用其解码数码 d_1 , 先解得

$$(M \cup T)^{a_1 d_1} \equiv (M \cup T) \pmod{m_1}$$

知甲提了自己的姓名(因 M 中提及甲的姓名), 于是知信息末端所剩的数码 T 为甲签字的电传数码。对此施用甲公布的发码组中的 a_0 作鉴定

$$T^{a_0} = N^{d_1 a_1} \equiv N \pmod{m_0} \quad (9)$$

$$\because a_0 d_0 \equiv 1 \pmod{m_0}$$

这样验证的确是甲的签字, 因为只有甲本人才会签 N^{d_1} , 其也可能用甲的发码组中的 a_0 解回得 N 。

研究问题

最近我们想到下面一个可能的电传签字冒签的问题。

问题：假设有人中途截获(或银行职员，或第三者)甲的一封电讯中，由于知甲的信中信尾的数字 $T(=N^d)$ 为甲名字 N 的签名，所以第三者可照用此数字就可仿签甲的电传签字了。

我们也有个方法来防止这种情形发生，但留待读者去思考，去解决这一有趣的问题。

参 考 资 料

1. M. R. Schroeder, "Number theory in science and communication", Springer-series in information science, Springer Verlag 1985.

第九章 电传打赌

§1 前 言

中国人剩余定理也就是当今中外数论书上对我们中国人老祖先的“韩信点兵术”或“孙子算经”中的“物不知数”的算法，及“九章算术”中的大衍求一算的基本理论。用现代术语就是求一次同余方程组解的主要理论。这个古老的数学计算法除了对数学其它分支上的应用外，较显著的是在今天最尖端的数学科技产品电脑中计算量理论，发挥了相当大的功效。在〔1〕中，我们也较完整地介绍了此一理论及它在电传公开密码上的应用。现我们再介绍此理论在一个性质相似的“电传打赌”的有趣应用。通常两个人(甲、乙)打赌，常喜欢用猜铜板正、反面来定输赢。即一个人掷铜板(或用手握住)由对方猜正、反。这种情形是两人面对面，看着掷出的是正、反面，所以没有什么可骗人的(但可耍赖，明明猜的是“正面”到时赖说猜的是“反面”)。如何把这种性质的“打赌”用电传通讯的方式(如打

电话、电报) 来进行呢? 首先我们把这种“打赌”, 用一个模式来表示:

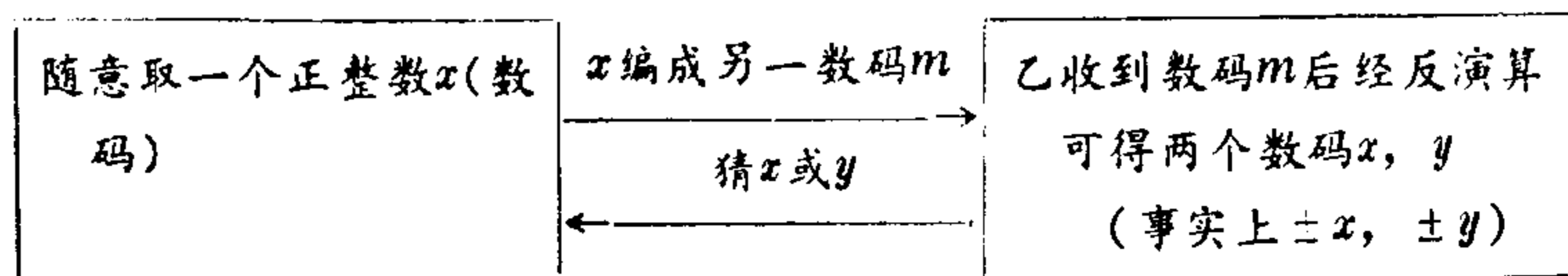


图11

当甲收到乙猜的 x 或 y 时, 特别当乙猜了 y , 甲告诉乙说“你输了”, 在没照相及第三者的旁证下, 怎可能心悦诚服地认了呢? 这是本文在下面要解说的, 为了方便读者的阅读, 我们先把要用到的类数理论(有关中国人剩余定理及其推广)作一扼要的介绍。

§ 2 有关的数学理论

2.1 一次同余方程

有两个正整数 a, b , 若 $a - b$ 为另一正整数 m 的倍数时, 我们称 a, b 对模 m 同余, 以 $a \equiv b \pmod{m}$ 表之, 这个概念在日常生活的计算中常常遇到(如中国人算生肖, 以12为模), 也为现代数论研究的一基本工具. 一个同余方程就是形如 $f(x) \equiv 0 \pmod{m}$ 的关系式, 其中 f 为一整数系数的多项式. 中国人定理是探讨当 $f(x)$ 为一次式的联立同余方程组解的存在性及解的表示, 在本文的应用中主要用到的是当 $f(x)$ 为二次式的二次同余方程解的结果。

我们先举一个由二个一次同余方程组成的方程组的理论作讨论, 中国人剩余定理就可顺理成章地叙述出来了。

定理9.1 方程组

$$\begin{cases} x \equiv a \pmod{m_1} \\ x \equiv b \pmod{m_2} \end{cases} \quad (2.1)$$

有解 x_0 的充要条件是 m_1 及 m_2 的最大公约数 (m_1, m_2) 可整除 $a-b$, 又设 x_0 为方程组的任意一个解, 则任何其他解(通解)可表为

$$x \equiv x_0 \pmod{[m_1, m_2]},$$

其中 $[m_1, m_2]$ 表 m_1 及 m_2 的最小公倍数.

证明: (充要性)

易见当 x_0 为方程组的解时其充要条件为存在有一整数 k , 使得下两式同时成立:

$$x_0 = a + km_1$$

$$\text{及 } a + km_1 \equiv b \pmod{m_2}$$

$$\text{或 } a - b \equiv km_1 \pmod{m_2}$$

而如此的 k 存在的充要条件为

$$(m_1, m_2) \mid (a - b) \quad (2.2)$$

(这可由若 m, n 两整数互质, 则必可找到两整数 α, β , 使得 $\alpha m + \beta n = 1$ 成立的事实得证).

现在我们看通解部分, 若 x_1 为另外解, 则

$$x_1 \equiv a \equiv x_0 \pmod{m_1}$$

$$x_1 \equiv b \equiv x_0 \pmod{m_2} \quad (2.3)$$

$\therefore x_1 - x_0$ 为 m_1, m_2 的一公倍数.

因此 $[m_1, m_2] \mid (x_1 - x_0)$

反之, 若

$$x_1 \equiv x_0 \pmod{[m_1, m_2]}$$

则很明显

$$x_1 \equiv x_0 \equiv a \pmod{m_1}$$

$$\text{及 } x_1 \equiv x_0 \equiv b \pmod{m_2}$$

即 x_1 亦为方程组的解。定理证毕。

同理不难把上面定理推广到任何多个同余方程组成的方程组而得到下面结果：

定理9.2 同余方程组

$$\begin{cases} x \equiv a_1 \pmod{m_1} \\ x \equiv a_2 \pmod{m_2} \\ \vdots \\ x \equiv a_n \pmod{m_n} \end{cases} \quad (2.4)$$

有解其充要条件为对任一对 $j, k (1 \leq j \leq n, 1 \leq k \leq n)$,

$$(m_j, m_k) \mid (a_j, a_k) \quad (2.5)$$

且若 x_0 为原方程组的任意一个解，则其任何一个解可表为

$$x \equiv x_0 \pmod{[m_1, m_2, m_3, \dots, m_n]}, \text{ 其中}$$

$[m_1, m_2, \dots, m_n]$ 表 $m_1, m_2, \dots, m_n$ 的最小公倍数。

当 $m_1, m_2, \dots, m_n$ 为两两互质的一组数时 ((即 $m_j, m_k = 1, j \neq k$)), 则由定理2知方程组(2.4)总有解。于是就得到下面的定理：

定理9.3 (中国人剩余定理)

设 $m_1, m_2, \dots, m_n$ 为一组两两互质的数，即 $(m_i, m_j) = 1$ ，当 $i \neq j$ 。令 $M = m_1 m_2 \dots m_n$ ，并设 $b_1, b_2, \dots, b_n$ 使得 y_j 分别为方程

$$y_j \frac{M}{m_j} \equiv 1 \pmod{m_j}, \quad j = 1, 2, \dots, n$$

y_j 的一个解。则方程组

$$\begin{cases} x \equiv a_1 \pmod{m_1} \\ x \equiv a_2 \pmod{m_2} \\ \vdots \\ x \equiv a_n \pmod{m_n} \end{cases}$$

的通解为

$$x \equiv a_1 b_1 \frac{M}{m_1} + a_2 b_2 \frac{M}{m_2} + \dots + a_n b_n \frac{M}{m_n} \pmod{M} \quad (2.6)$$

例9.1 试解同余方程组

$$\begin{cases} x \equiv 5 \pmod{7} \\ x \equiv 7 \pmod{19} \end{cases} \quad (2.1)$$

解: 此时 $m_1 = 7$, $m_2 = 19$ 及 $M = m_1 \times m_2 = 133$

b_1 为方程

$$y \frac{M}{7} \equiv 1 \pmod{7}$$

即 $19y \equiv 1 \pmod{7}$

的一个解 y , 不妨取 $y = 3 = b_1$

同理可得 $b_2 = -8$ 为方程

$$y \frac{M}{19} \equiv 1 \pmod{19}$$

之一解。

于是原方程组的通解为

$$x \equiv 5 \times 57 + 8 \times (-56) \equiv 103 \pmod{133}$$

2.2 高次同余方程

我们将讨论当 $f(x)$ 为整数系数多项式时同余方程

$$f(x) \equiv 0 \pmod{m} \quad (2.7)$$

的解。

由于若 a 为一满足方程(2.6)的一整数解, 则任何与 a 对 m 同余的数 b (即 $b \equiv a \pmod{m}$) 亦同时为(2.7)的解。这可由若 $f(a) \equiv 0 \pmod{m}$, 则对任何整数 $f(a + tm) \equiv 0 \pmod{m}$, 得证。所以一般我们要知道方程(2.6)有多少解, 是指有多少互不同余的解。我们将证明一般若 m 不是某个质数的幂次的话, 则一般解

方程 $f(x) \equiv 0 \pmod{m}$ 可化为解一组模数变小了的同余方程.特别当 $f(x)$ 可分解成线性因子时,则就归化到解一组一次同余方程,就可利用中国人剩余定理把通解求出来了.

定理9.4 设 $f(x)$ 为整系数多项式

设 $M = m_1 \cdot m_2 \cdots m_n$, 其中 $m_1, m_2, \cdots, m_n$ 为两两互质的正整数, 则整 a 为方程

$$f(x) \equiv 0 \pmod{M}$$

的解, 其充要条件是 a 为下列同余方程组之解:

$$\begin{cases} f(x) \equiv 0 \pmod{m_1} \\ f(x) \equiv 0 \pmod{m_2} \\ \vdots \\ f(x) \equiv 0 \pmod{m_n} \end{cases}$$

证: 必要性

若 $f(a) \equiv 0 \pmod{M}$, 则必然有

$$f(a) \equiv 0 \pmod{m_i}, \quad i = 1, 2, \cdots, n.$$

充分性: 反过来设 a 满足

$$f(x) \equiv 0 \pmod{m_i}, \quad i = 1, 2, \cdots, n.$$

则 $f(a)$ 满足下列方程组

$$\begin{cases} y \equiv 0 \pmod{m_1} \\ y \equiv 0 \pmod{m_2} \\ \vdots \\ y \equiv 0 \pmod{m_n} \end{cases}$$

于是由定理3(或2.7)可得知

$$f(a) \equiv 0 \pmod{m_1 \cdot m_2 \cdots m_n}$$

亦即 a 为 $f(x) \equiv 0 \pmod{M}$

的一解.

例9.2 设 $f(x) = x^5 - 3x^4 + 7x^3 - 2x^2 - 9x + 6$. 试解

$$f(x) \equiv 0 \pmod{165}.$$

解: $\because 165 = 11 \times 5 \times 3$, 所以我们来解

$$\begin{cases} f(x) \equiv 0 \pmod{11} \end{cases} \quad (2.8)$$

$$\begin{cases} f(x) \equiv 0 \pmod{5} \end{cases} \quad (2.9)$$

$$\begin{cases} f(x) \equiv 0 \pmod{3} \end{cases} \quad (2.10)$$

对方程(2.8)而言, 由于模数为11, 故我们只需考虑试 $x = 0, 1, 2, 3, \dots, 10$, 看 $f(x) \equiv 0 \pmod{11}$ 是否成立? 得满足此关系式的是 $x = 1, x = 6$ 及 $x = 8$, 因此(2.8)的通解为

$$\begin{cases} x \equiv 1 \pmod{11} \\ x \equiv 6 \pmod{11} \\ x \equiv 8 \pmod{11} \end{cases} \quad (A1)$$

同理得 $f(x) \equiv 0 \pmod{5}$ 的通解为

$$\begin{cases} x \equiv 1 \pmod{5} \\ x \equiv 2 \pmod{5} \\ x \equiv 3 \pmod{5} \end{cases} \quad (A2)$$

及 $f(x) \equiv 0 \pmod{3}$ 的通解为

$$\begin{cases} x \equiv 0 \pmod{3} \\ x \equiv 1 \pmod{3} \end{cases} \quad (A3)$$

所以由通解(A1)、(A2)、(A3)的组合一共可得 $3 \times 3 \times 2$ 共18组解. 如取任一组, 例如

$$\begin{cases} x \equiv 1 \pmod{11} \\ x \equiv 1 \pmod{5} \\ x \equiv 0 \pmod{3} \end{cases}$$

由定理9.3可得一解为 $x = 111$, 因此解同时满足(A1—3), 故为 $f(x) \equiv 0 \pmod{165}$ 之一解.

下面我们来讨论如何判定方程 $f(x) \equiv 0 \pmod{m}$, m 为质数, 最多有几个不同余的解? 这一结果在讨论 $m = p \times q$ 两个不等的

质数，也是在我们电传打赌将会用到的情形。我们先证明一个一般的结果。

定理9.5 (因式定理) 设 $f(x)$ 为一非常数的 n 次整系数多项式，则整数 a 为方程

$$f(x) \equiv 0 \pmod{m}$$

的充要条件为存在一次数 $n-1$ 的整系数多项式 $g(x)$ 使得

$$f(x) \equiv (x-a)g(x) \pmod{m}.$$

证：充分性

$$\text{若 } f(x) \equiv (x-a)g(x) \pmod{m}.$$

$$\text{则 } f(a) \equiv (a-a)g(a) \equiv 0 \pmod{m}.$$

必要性

我们用长除法将 $f(x)$ 被 $x-a$ 除，则可得一商式 $g(x)$ 及余式 r (为一常数)，很明显 $g(x)$ 为整数系数多项式，次数为 $n-1$ ，及 r 为一整数。于是

$$f(x) = (x-a)g(x) + r$$

当 $x=a$ 时

$$f(a) = (a-a)g(a) + r \equiv r \pmod{m}$$

现因 a 为 $f(x) \equiv 0 \pmod{m}$ 之解，所以 $f(a) \equiv 0 \pmod{m}$ 。于是由上式得

$$r \equiv 0 \pmod{m}$$

因而

$$f(x) \equiv (x-a)g(x) \pmod{m}$$

定理证毕。

特别当 m 为一质数时，我们不难得到下面定理：

定理9.6 设 P 为一质数， $f(x)$ 为一 n 次整数系数多项式，则若 $a_1, a_2, \dots, a_t$ 为 t 个属于方程 $f(x) \equiv 0 \pmod{p}$ 的解，则必存在有一个 $n-t$ 次的多项式 $g(x)$ ，使得

$$f(x) \equiv (x - a_1)(x - a_2) \cdots (x - a_t)g(x) \pmod{p}$$

由此立即可得以下推论：

定理9.7 设 p 为一质数， $f(x)$ 为一 n 次整数系数的多项式。则 $f(x) \equiv 0 \pmod{p}$ 至多有 n 个互不同余的解。

§3 二次同余方程及二次残余的一些性质

我们在电传打赌中要用到的是二次同余方程及二次残余的一些性质。

定义9.1 我们称一个整数 a 为一正整数 m 的二次残余是指 $(a, m) = 1$ 及同余方程 $x^2 \equiv a \pmod{m}$ 有解。若 $x^2 \equiv a \pmod{m}$ 无解，则称 a 为 m 的二次非残余。

例如，对 $m = 11$ 时 $a = 1, 3, 4, 5$ 及 9 为其二次残余，但 $2, 6, 7, 8$ 及 10 就是 m 的二次非残余。

一般可以证明对任何一奇质数 p ，在 $1, 2, 3, \cdots, p-1$ 共 $p-1$ 个数中，一半是 p 的二次残余，一半是 p 的二次非残余。这个事实可由下面结果导出：

定理9.8 设 p 为一奇质数， a 为一不被 p 整除的整数。则方程 $x^2 \equiv a \pmod{p}$ 或无解，或有两个以 p 为模但不同余的解。

证：若 $x^2 \equiv a \pmod{p}$ 有一解，设其为 x_0 ，则我们不难验证 $x = -x_0$ 为第二个与 x_0 不同余的解（ $\because$ 若 $x_0 \equiv -x_0 \pmod{p}$ ，则 $2x_0 \equiv 0 \pmod{p} \Rightarrow p$ 为一偶数）。

现在我们证至多有两个以 p 为模互不为同余的解。因若 $x = x_0$ 及 $x = x_1$ 皆为 $x^2 \equiv a \pmod{p}$ 之解，则依据定理9.7知 $x^2 \equiv a \pmod{p}$ 不可能有其它不同余的解，在我们判定什么情形下的 a 及 p ，方程 $x^2 \equiv a \pmod{p}$ 有解前，先引进下面一常用的记号。

定义9.2 设 p 为一奇质数， a 为一不被 p 整除的整数，则勒兼德符号： $\left(\frac{a}{p}\right)$ 定义如下：

$$\left(\frac{a}{p}\right) = \begin{cases} 1 & \text{若 } a \text{ 为 } p \text{ 的一二次残余,} \\ -1 & \text{若 } a \text{ 非为 } p \text{ 的一二次残余.} \end{cases}$$

换言之， $\left(\frac{a}{p}\right) = 1$ 表示 $x^2 \equiv a \pmod{p}$ 有解，而 $\left(\frac{a}{p}\right) = -1$ 表示 $x^2 \equiv a \pmod{p}$ 无解。

现我们可以介绍 $x^2 \equiv a \pmod{p}$ 有解或无解的检定。

定理9.9 (尤拉检定) 设 p 为一奇质数， a 为一不被 p 整除之整数，则

$$\left(\frac{a}{p}\right) \equiv a^{(p-1)/2} \pmod{p} \quad (2.11)$$

证：假设

$$\left(\frac{a}{p}\right) = 1$$

则此表示方程 $x^2 \equiv a \pmod{p}$ 有解，不妨设其为 $x = x_0$ ，则由费马小定理可得

$$a^{(p-1)/2} = (x_0^2)^{(p-1)/2} = x_0^{p-1} \equiv 1 \pmod{p}$$

这也就是说 若 $\left(\frac{a}{p}\right) = 1$ 则

$$\left(\frac{a}{p}\right) \equiv a^{(p-1)/2} \pmod{p} \text{ 成立.}$$

现看 $\left(\frac{a}{p}\right) = -1$ 时如何？

这表示 $x^2 \equiv a \pmod{p}$ 无解。由此不难知道对任一整数 j 其满足 $1 \leq j \leq p-1$ ，则必存在一唯一的整数 i ， $1 \leq i \leq p-1$ 使得 $ij \equiv a \pmod{p}$ 成立。因此我们可以把 $1, 2, \dots, p-1$ 两个如此配起来，最后一齐相乘得

$$(p-1)! \equiv a^{(p-1)/2} \pmod{p}$$

由Wilson (威尔逊定理) $(p-1)! \equiv -1 \pmod{p}$. 此与上式结合得

$$\left(\frac{a}{p}\right) = -1 \equiv a^{(p-1)/2} \pmod{p}$$

因而我们证明了不论 $\left(\frac{a}{p}\right)$ 为 1 或 -1, 总满足(2.11).

注: 有关勒兼德记号的另一重要性质是

$$\left(\frac{a}{p}\right)\left(\frac{b}{p}\right) = \left(\frac{ab}{p}\right) \quad (2.12)$$

这是说若 $x^2 \equiv a \pmod{p \times q}$ 有解, 则 a 为 p 及 q 的二次残余.

电传打赌中特别要求用到的是取 $p \equiv q \equiv 3 \pmod{4}$ 的情形.

定理9.10 设 p, q 为质数, 满足 $p \equiv q \equiv 3 \pmod{4}$, 且 $x^2 \equiv a \pmod{p}$ 及 $x^2 \equiv a \pmod{q}$ 有解, 则

$$x \equiv \pm a^{(p+1)/4} \pmod{p} \quad (2.13)$$

及

$$x \equiv \pm a^{(q+1)/4} \pmod{q} \quad (2.14)$$

分别为上述两方程之解.

证: 因由假设 $x^2 \equiv a \pmod{p}$ 及 $x^2 \equiv a \pmod{q}$ 有解, 所以由定理9.9可得

$$a^{(p-1)/2} \equiv \left(\frac{a}{p}\right) = 1$$

及

$$a^{(q-1)/2} \equiv \left(\frac{a}{q}\right) = 1$$

因此

$$(\pm a^{(p+1)/4})^2 = a^{(p+1)/2} = a \cdot a^{(p-1)/2} \equiv a \pmod{p}$$

$$(\pm a^{(q+1)/4})^2 = a^{(q+1)/2} = a \cdot a^{(q-1)/2} \equiv a \pmod{q}$$

验证了(2.13)及(2.14)分别为 $x^2 \equiv a \pmod{p}$ 及 $x^2 \equiv a \pmod{q}$ 的解。

今 $a^{(p+1)/4}$ 与 $-a^{(p+1)/4}$ 及 $a^{(q+1)/4}$ 与 $-a^{(q+1)/4}$ 分别为方程 $x^2 \equiv a \pmod{p}$ 及 $x^2 \equiv a \pmod{q}$ 两互不同余的解。所以由中国人剩余定理(定理9.3)可以把 $x^2 \equiv a \pmod{p \times q}$, a 与 p, q 互质, p, q 为奇质数的4个不同余的解得出。我们可将以上分析叙述成一定理如下:

定理9.11 设 $n = p \times q$, p, q 为两不等的奇数并假定

$$x^2 \equiv a \pmod{n}, \quad 0 < a < n \text{ 有一解, } x = x_0.$$

则 $x^2 \equiv a \pmod{n}$ 4个不同余的解可表为

$$\begin{array}{ll} \text{(I)} \quad x \equiv x_1 \pmod{p} & \text{(III)} \quad x \equiv p - x_1 \pmod{p} \\ & x \equiv x_2 \pmod{q} \\ \text{(II)} \quad x \equiv x_1 \pmod{p} & \text{(IV)} \quad x \equiv p - x_1 \pmod{p} \\ & x \equiv q - x_2 \pmod{q} \end{array}$$

其中 x_1 及 x_2 分别满足 $x_0 \equiv x_1 \pmod{p}$, $0 < x_1 < p$ 及 $x_0 \equiv x_2 \pmod{q}$, $0 < x_2 < q$.

注:若我们将解(I)及(II)分别以 x 及 y 记之, 则解(III)及(IV)为 $n - y$ 及 $n - x$, 则 $x + y \equiv 2x_1 \not\equiv 0 \pmod{p}$ 及 $x + y \equiv 0 \pmod{q}$, 这样就得

$$(x + y, n) = q \quad (2.15)$$

同理可得

$$(x + (n - y), n) = p \quad (2.16)$$

或 $(x - y, n) = p$

这儿记号 (a, b) 表 a, b 两数的最大公约数。

§ 4 打赌步骤

有了以上的准备，我们可把电传打赌的数学模式及步骤叙述如下：

假定有甲、乙两方要借电传（电话、电报…）打赌，并约定由甲选一数字，而由乙来猜。

1. 乙先选两个很大而不等的质数 p 及 q 满足 $p \equiv q \equiv 3 \pmod{4}$ （注意形如 $4k+3$ 的质数有无穷多的！），乙将 p 、 q 的乘积 n （即 $n = p \times q$ ）先拍送给甲方。

2. 甲收到 n 后任选一个比 n 小但与 n 互质的任一个正整数 x 然后计算 $x^2 \equiv a \pmod{n}$ 使得 $0 < a < n$ （即 a 为 x^2 对模 n 的最小正残余：如 $n = 3 \times 7 = 21$ ， $x = 2$ 则 $x^2 \equiv 4 \pmod{21}$ 于是 $a = 4$ ）注意：这里的 x 就是甲所选的数字！

3. 甲将 a 拍送给乙（准备由乙来猜了）

4. 乙收到 a 后，解方程 $w^2 \equiv a \pmod{pq}$ ，（这个方程是可解的）因步骤二中，甲选的 x 就满足此一方程！且此方程的解可由定理9.10而得。再依据定理9.11及其注可得解 $w = x$ 或 $n - x$ 及 $w = y$ 或 $n - y$ ，我们也可看有解 $w = \pm x$ 及 $w = \pm y$ 。

5. 乙选 $\pm x$ 或 $\pm y$ 中之一拍送给甲（这是乙猜的结果！）

6. (I) 若甲收到的是 x （或 $-x$ ），则具有回答乙说他猜对了。

(II) 若甲收到的是 y （或 $-y$ ），则告诉乙说他猜错了。

对于(I)的情况如果乙存心想赢，他不会再去怀疑甲在骗他，但如(II)发生而乙又不确定甲是有所依据而讲的，这是在骗他，怎么办？

下一步骤是由乙来检验甲是否是诈骗才需要的。（也即乙

拍送给甲的是 $\pm y$ ，而甲所选的是 x ，乙怎样相信甲没骗他?)

7. 乙要求甲把 n 作一分解。

这时如甲没骗乙，则甲自己握有的数据是 x ，加上乙拍送的 y (或 $-y$) 则由定理9.11后之注中的式(2.15)或(2.16) 可把 p 或 q 得出，因而就可把 n 的因子 p 、 q 都找出来。所以当甲把 p 、 q 拍送给乙，乙就只有认输了。因我们知道给定一个很大的整数 n ，且知道 n 为某两个质数的相乘积，但将此两因子求出一般而言是很费时的。这也是电传打赌公平保证。

我们现举一个计算的例子

例9.3 假定我们事先知道方程 $x^2 \equiv 860(11021)$ 有解，如

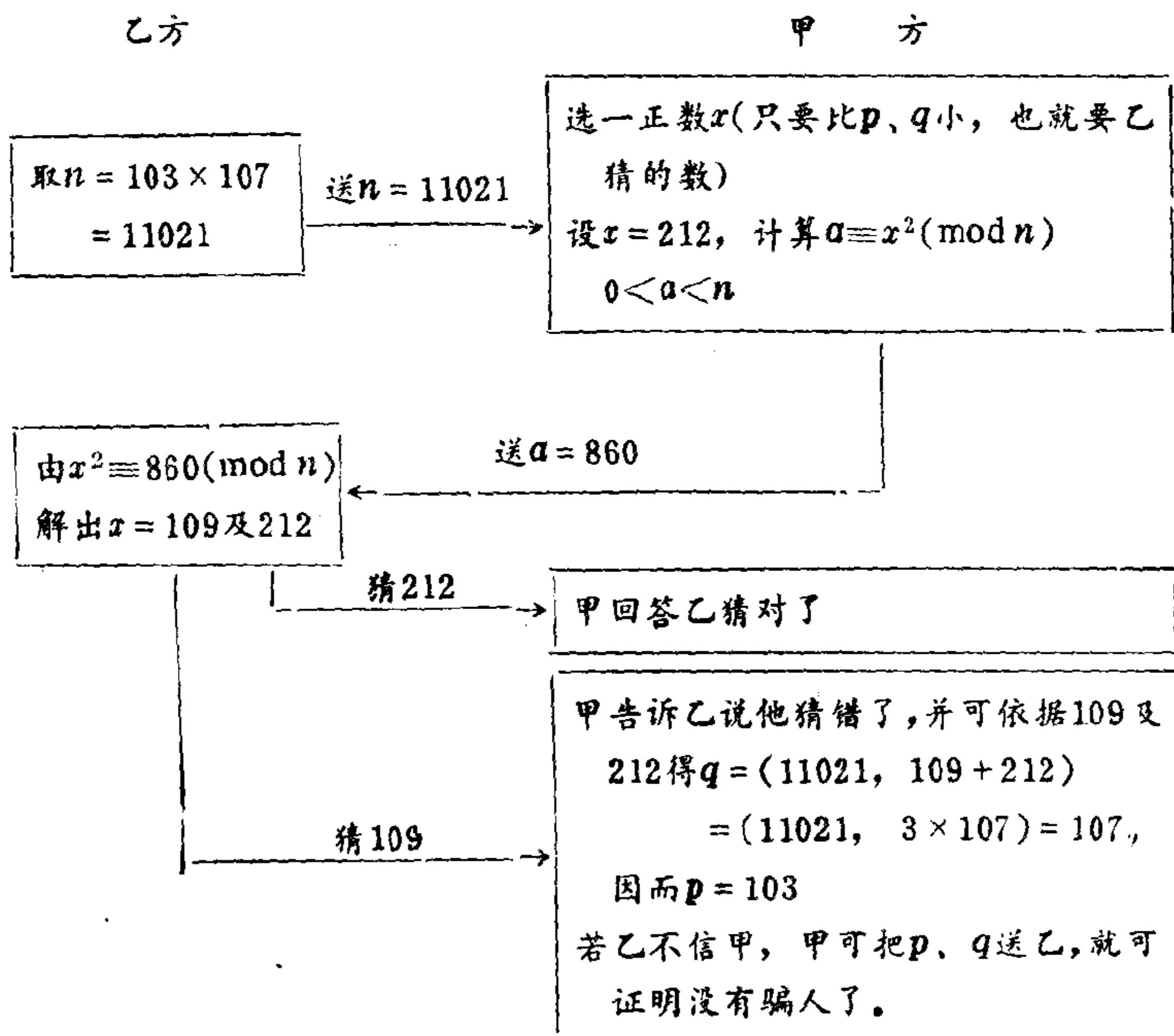


图12

何来解此方程呢?

解: 现 $11021=103\times 107$ (两个奇质数之乘积),

且 $103\equiv 107\equiv 3(\bmod 4)$ 依据定理10, 知

$$x^2\equiv 860\equiv 36(\bmod 103) \text{ 及}$$

$$x^2\equiv 860\equiv 4(\bmod 107) \text{ 的4个不同余的解分别为}$$

$$x\equiv \pm 36^{(103+1)/4}\equiv \pm 36^{26}\equiv \pm 6(\bmod 103) \text{ 及}$$

$$x\equiv \pm 4^{(107+1)/4}\equiv \pm 4^{27}\equiv \pm 2(\bmod 107).$$

由此两联立方程组再引用中国人剩余定理可得原方程的解为

$$x\equiv \pm 212, \pm 109(\bmod 11021)$$

$$\text{检验 } (\pm 212)^2\equiv (\pm 109)^2\equiv 860(\bmod 11021)$$

利用上例我们可举一个打赌的实例(见图12)。

§5 研 讨

(I) 这样的打赌有无失公平性呢?由前面的步骤我们可看到甲有两个平等的机会 ($\pm x$ 或 $\pm y$), 使得他能分解 n (收到的是 $\pm y$) 或使得他不能分解 n (收到的是 $\pm x$), 乙有两个公平的机会去猜中或没猜中, 甲原先定的 x , 所以该算是公平的。

(II) 假如甲存心想输, 乙有无法防止?

也就是说明明乙拍送给甲的是 $\pm y$, 而不是 $\pm x$, 但甲可装收到的是 $\pm x$, 而告诉乙说他赢了怎么办?

参 考 书 目

1. 杨重骏、杨照崑 “整数论及其应用” 华东书局, 1982.
2. Kenneth Rosen, Elementary numebr theory and its appliations, Addison-Wesley Pudlishing Co. 1984.

第十章 统计学与模拟实验

§ 1 统计学概说

我们日常生活中所作的决定都是根据我们的知识与经验，严格地说，知识也是前人的经验，因此我们可以概括的说，人类之所以为万物之灵，在于人类能从经验中吸取教训，进而对以后处理类似事件时的改善。生物在进化过程中当然也有因环境而改变其生存方式的能力，但这种能力是由于基因在遗传上的改变，所以要历时很长久，而人类，因有文字与语言，可以把经验在当代就记录下来，迅速的传开，有此几十代前人认为是新奇的经验，现在已成为常识。像地球是圆的，某些疾病可以传染，雷电不是雷公、雷母用以惩罚不孝的人，这些都是前人经历过来的事实，我们不必一一再亲自去经历。但前人有时也会传下来一些不正确的结论，像“腐草化萤”、“乌鸦反哺尽孝”等等，现已经经过仔细的观察及研究之后而认为是不正确的了。由于科学进展和知识领域日新月异，我们不能只停留在现

有书本上的知识，还得靠一些书本上没有的——或是亲身经历，或是人们的故事，甚至道听途说——来帮助我们对一些新的事情做决定。

“经验是最好的老师，只是学费太贵”，如果我们想把我们的经验告诉社会大众，使一般人不必再重复去经历而可得到我们所经历的教训，那我们就必须小心，避免以自己的经历而造成了别人的偏见，结果对另一些人而言，不但无益，反而有害。很多固执的长者往往犯了这个毛病，他们以其自身的经验作为真理，强迫他们的子女或下属去服从。比如说，不久以前有一位医学教授声称“吃鸡蛋不会引起胆固醇过多症”，他的理由是，因为他每天吃一个鸡蛋，而至今血压不高，血中的胆固醇正常。我们又常听到某些偏方，像草根、树皮，甚至香灰可以治疗癌症，说的人未必有意骗人，他也许是根据他自身所亲见的一个事实，但往往因此误了另一个病人用正确方法治病的时机。有人相信梦、有人相信手相，这些事未必没有灵验的时候，但把一个有灵验的例子当作一般的事实来传播，常常会害人不浅。因此当我们在宣布一项新发现的时候，要特别慎重。为了防止以一概全的错误，我们应有一套比较客观的方法来认识及解决周围发生的现象，这个方法就是统计学。

统计学的主要目的在从一些样本中，推论某个假定正确的可能性。比如说，有一个人看见某个小孩子天天看电视，但功课却很好，因此他假说“小孩子成绩好”。为了证实这个假说，他从某国随便挑选了二十个功课好的小朋友及二十个功课不好的小朋友，这就是此国小学生中的一个样本。结果他发现这二十个功课好的小朋友平均比那二十个功课不好的小朋友电视看得多。那么现在他是不是可以下结论说“多看电视的小孩子成绩好”呢？这中间可有两个问题讨论：

第一，他所取的样本的大小。四十个小孩子够不够大？我们不妨这样想，若他只取了一个功课好的小孩及一个功课不好的小孩，他有没有资格下上面所说的结论？你一定说没有。那么各二个呢？还没有？这样一直下去，他抽了多少个样本才可以下一个结论呢？很显然的，这个问题的答案不可能是一个数目。比如说答案是在二十个样本以上就可以下结论，那么十九与二十之间为什么有这样大的一个鸿沟呢？对这个问题的一个正确的答案是，有多少样本，你就只可以下有多少把握的结论。样本愈大，就对结论正确性的把握也就愈大，这个把握是以或然率来衡量的。当你取一个样本的时候，你也可以根据样本的结果来下一个结论，但你结论有50%不正确的或然率，因此你不便发表，但若你的样本够大，你的结论有99%正确的或然率，则你不妨发表，并诚实的告诉别人的你结论可能错误的或然率是1%。

如何计算这个或然率往往不是一个简单的问题。古典统计学都是从解析理论来寻求计算公式。但由于计算机快速的发展，或然事件往往很容易由计算机来模拟(参看本书第七章第3节)，我们将以统计学中的假说检定来看模拟实验的应用。

§ 2 假说检定

假说检定 (hypothesis testing) 在正确表明某假说为正确之可能性。比如说有人提出一个假说：

假说：用右手写字的人左眼视力比右眼好。

为了检定此一假说的正确性，我们预备在某大学中随意取出100个用右手写字的大学生，比较他们两眼的视力。表1显示

了四种可能得到的抽样结果。

表1 表中数字表示在100人中左眼视力较好的人数

	可能性之一	可能性之二	可能性之三	可能性之四
结 果	70	52	50	42
结 论	?	?	?	?

如果有这样的结果，我们的结论是什么？在直觉上，我们很可能认为若我们抽样的结果是第一、二种情形，则可以认为假说为真，而在第三、四种情形下认为假说为错误（第三种情形在直觉上是眼力相等而非较好）。但如果我们举出了这样的报告，有没有错误的可能性？当然有。假定该大学有三万个用右手写字的学生而其中有二万个是右眼比左眼强，而另一万个是左眼比右眼好。这时假说可以说是不成立的，但在100人抽样中会不会产生表1中的各种可能性呢？都有可能。因此虽然假说不成立，我们会有可能得到一、二的情形而导至一个错误的结论。因此我们知道在这种抽样的情形下，我们对假说的肯定或否定都不可能100%的正确。就是我们抽样大到1000人，甚至10000人，我们仍然得不到100%正确的结论。太大的抽样往往花费很大，耗时很久。如果我们的能力只能作100人的抽样，那么表1的各种情形对我们检定此假说有无用处？答案是：

“表1对作一个100%正确的假说检定可能无用，但它有时可以帮助我们对假说做一个相当可靠的检定。”

这种可靠性的计算法我们立刻就会谈到。但我们先把它它们列在表1的第三行中。

从表2中我们可以看出我们对情形之一所下的结论很有把握，因为此结论错误的或然率只有万分之一。但在情形二、三

表2

在100个抽样中左眼视力较好的可能人数

	情形之一	情形之二	情形之三	情形之四
人 数	70	52	50	42
暂取结论	假说正确	假说正确	假说不对	假说不对
暂取结论不正确时， 会发生这类样本的或 然率	0.0001	0.4013	0.50	0.0548

时所下的结论就非常危险，因为错误的可能性太高了，也就是如果结论是如此得来的，则常常会错。我们自然不应该在这种情形下草率地宣布我们已知道此假说的正确性。至于在第四种情形，因犯错误的或然率不大不小，百分之五左右，我们就很难决定此假说的正确性。一般的做法是宣布结论并告诉读者此结论可能错误的可能性，让读者心理有一个准备。

要看表2中几个或然率是怎么算出来的，我们先从简单的例子想起。假定一个袋子里有黑球与白球，且数目很大，但黑白之比例我们不知道。设假说是：

假说：袋中黑球多于白球。

为了检定此假说的正确性，我们将取出 n 个球作观察。先令 $n=1$ ，即取出一球，假定此球为黑色。若我们据此声称假说为正确，那么我们的声称发生错误的或然率是多少？令

p = 袋中黑球之比例。

则取一球而得一黑球之或然率为 p ，若 $p \leq 0.5$ ，则我们的结论发生错误，因此“在假说不正确的情形下，我们有 p 的或然率得到一个这样的样本而下了一个错误的结论”。 p 是未知数，但可高达0.5，故我们说我们的结论错误的或然率可以高达0.5（但不会高过0.5，因 $p > 0.5$ 时，假说就正确了）。“结论错误的或然率”

这个词句可能引起误会，它不是指这次结论发生错误的或然率”因为袋中的球是一定的，假说对不对也是一定的。因此声称“假说正确”这件事非对即错，没有或然率而言。“结论错误的或然率可以高达0.5”这句话的意思是：

我们若以目前这种态度下结论（即取一球，得一黑球则宣布袋中黑球多于白球），我们所有的结论中，可以高达一半是错误的。

再把这个问题换一个角度看。假定我们得到了一个白球，我们仍宣称假说“袋中黑球多于白球”正确，我们结论错误的或然率是多少？在这种情形下，我们得到这样样本的或然率是 $1-p$ ，在 $p \leq 0.5$ 的情形下， $1-p$ 可以高达1。因此我们结论错误的或然率可以高达1。我们再强调一遍，这并不是这次决定错误的或然率是1或高达1，因这次对不对已经定了，而是：

我们若以目前这种态度下结论（即取一球，看到了白球，仍宣布黑球多的假说成立），我们所有的结论可以全是错的。

这是很显然的。因为这种几乎是不看事实下结论（看到白球仍说黑球多），自然可能全错。

现在假定 $n=2$ ，即取2球。因袋中的球极多，我们可以用二项展开式来计算取出球各种情形的或然率。现假定取出的两个球都是黑球，如果因此我们认定假说成立，则可能发生错误的或然率 α 是在 $p \leq 0.5$ 时得两个黑球的或然率，即

$$\alpha = p \text{ (在二样本中得两个黑球} | p \leq 0.5 \text{)}.$$

在上式中的“|”号表示已知右式 $p \leq 0.5$ 成立。若再以 X 表示2球中黑球的数目，则上式成为

$$\alpha = p(X=2 | p \leq 0.5) = p^2; \quad p \leq 0.5$$

上式的 α 可以达到 $0.5^2 = 0.25$ ，即

$$\alpha = p^2 \leq 0.25$$

且可以等于0.25。也就是说在这种情形下宣布假说 $p > 0.5$ 成立造成错误的或然率不会大于0.25，但可以高达0.25。

再往下看，若任意取10个球全都是黑球，那么我们可能错误的或然率就变成：

$$\alpha = p(X = 10 | p \leq 0.5)$$

依前法可得

$$\alpha = p^{10} \leq 0.5^{10} = 0.000976 \approx 0.001 \quad (1)$$

即在这种情形下宣布假说 $p > 0.5$ 成立发生错误之可能性不大于0.001，至多也只有0.001，即千分之一。我们应可以宣布 $p > 0.5$ 成立。

现在还剩下一个问题，即如果10个球中是9黑1白，我们若结论假说 $p > 0.5$ 成立，则此结论错误之或然率是多少？这就触及统计学校深一层的理论。前面谈过当我们说“我们宣布 $p > 0.5$ 成立，而可能错误之或然率为0.001”这句话时，其实它的含义与字面稍有不同。我们再强调一次这个常常引起误解的观念：当我们宣布 $p > 0.5$ 成立时，我们是否正确或错误已经定了（但我们不知道），因袋中的黑球与白球是一定的，是不是 $p > 0.5$ 也是一定的。因此我们可能错误的或然率是0.001是指：在一千次这样的决定中，可能有一次是错的。因此当我们在做决定时，我们所持的不单是这一次的决定，而是指所有其他类似这样的决定。因此我们如果在10个球中有9黑1白就宣布假说 $p > 0.5$ 成立，则下次如果在10球中有10黑0白，我们一定也会宣布假说 $p > 0.5$ 成立。因此我们在9黑1白宣布假说 $p > 0.5$ 成立，其实是指：若在10球中，黑球多于或等于9个，宣布假说成立。因此在此情形下可能造成错误的或然率是

$$\alpha = p(X \geq 9 | p \leq 0.5), X \sim b(10, p)$$

$$\begin{aligned}
 &= p^{10} + \binom{10}{1} p^9 (1-p)^1 = p^9 (p + 10(1-p)) \\
 &= p^9 (10 - 9p)
 \end{aligned} \tag{2}$$

很显然由微分可知 α 之极点在 $p=0,1$,故在 $p \leq 0.5$ 时 α 之极大值为

$$\alpha \leq 0.5^9 (10 - 9 \times 0.5) = 0.011$$

故在此情形下最大错误的或然率也不过是百分之一。是否接受此假说,则要看此结论之重要性,若此假说与千万人生命有关,则此错误之或然率仍嫌太大。故此在发表论文时,一般著者都附上此假说是在何种可能错误的或然率下接受的,因为使用同一结果的别人,可能需要不同程度的保证。

再往下看,若在10球中所得为8黑2白,则同理可得

$$\begin{aligned}
 \alpha &= p(X \geq 8 | p \leq 0.5), X \sim b(10, p) \\
 &= p^{10} + \binom{10}{1} p^9 (1-p) + \binom{10}{2} p^8 (1-p)^2 \\
 &\leq (1 + 10 + 45) \times 0.5^{10} = 0.055
 \end{aligned} \tag{3}$$

且可高达0.055。

现在可以回到表2中的或然率了。我们且取第一种情形为例。设 X 为100个人中左眼较好的人数, p 为左眼较好者在母体中的比例,则在第一种情形下接受 $p > 0.5$ 发生错误之或然率为

$$\begin{aligned}
 \alpha &= p(X \geq 70 | p \leq 0.5) \\
 &\leq p(X \geq 70 | p = 0.5) \\
 &= \sum_{i=70}^{100} \binom{100}{i} (0.5)^{100} \\
 &\leq 0.0001
 \end{aligned} \tag{4}$$

表2中其他几个或然率之值也可依法求出。

现在我们可以看到假说检定的程序是

1. 你想证实的假说 (一般以 H_0 表示), 和当你假说不对时,

我们以 H_0 表示²。对黑白球的例子而言，

H_0 : 袋中黑球比白球多 ($p > 0.5$)

H_0 : $p \leq 0.5$

2. 取一个合理的判别法。在黑白球问题中，如果抽样中黑球太多时，即可宣称 H_0 为正确。什么是太多则要由 α 来决定，即求出：

3. 若在目前的情形下宣称 H_0 成立时所可能错误的或然率（也可称为所冒之风险，通常以 α 表示³）。

读者可以看出前面（1—4）等式全是由上面三个程序所导引出来的。

但要求出 α 并不都像（1—4）式那样容易，今再举一例。设某种汽车有六种颜色，白，红，橙，黄，灰，蓝，今在某地各色之销售量分别为15，5，6，16，9，9辆，试问我们可否确定人们对各色需要的比例不同。依假说检定之程序：

1. H_0 : 人们对各色之需求比例相同。

H_0 : 人们对各色之需求比例不相同。

2. 我们取 H_0 若各色之要求量相差太大。以公式表示的方法之一为⁴：

$$Q = \sum_{i=1}^6 |o_i - n/6| > c \quad (5)$$

式中 o_i 为各色之销售量， n 为全部售量60， $n/6 = 10$ 表示若在 H_0 的情形下 o_i 最可能之值。因此若 o_i 与10相差太多，则 H_0 不会成立了。

3. 现在要求在 H_0 的情形，

$$p \left(\sum_{i=1}^6 |o_i - 10| \geq c \right) \quad (6)$$

而式中 c 为现今之 Q 值即 22。若用解析数学，(6) 非常不容易计

算。但若用模拟实验，在 H_0 成立的情形下，反复取60球看 Q 之值有多少机会大于22。我们实际上取了1000次模拟取样实验，所得到的结果是

$$\alpha = P, \left\{ \sum_{i=1}^6 |o_i - 10| \geq 22 \mid H_0 \right\} \doteq 0.05$$

因此种 α 位于临界点，最好是将结果与 α 同时公布给使用者。

现在回过来谈一下模拟取样的方法。因在 H_0 成立的情形下，各色汽车需求相同，因此我们每取一辆汽车，分别以 $\frac{1}{6}$ 的机会令它为六色中之一种。那么在60辆汽车抽完之后，我们会有一个“在各色需求相同情形下（即 H_0 ）”的一个可能发生的样本。若结果是10, 12, 14, 16, 3, 5，则 $Q=24$ 。也就是说在 H_0 成立是 Q 有可能是24，但可能性的大小则必须再多抽样。我们如此做了一千次，发现有52次 $Q \geq 22$ 。因此我们知道 $P, \{Q \geq 22\}$ 的机会约为千分之52，即0.05。

再举一个例子，设有某种新教学法，想与传统的方法相比较。现对二套教学法各取了四个程度相当的班级。在用同样的时间材料教授之后，做一次会考，成绩如下表：

表3

新方法(4班)	72(x_1), 64(x_2), 68(x_3), 76(x_4); 平均($\bar{x}$) = 70
旧方法(4班)	60(y_1), 64(y_2), 70(y_3), 78(y_4); 平均($\bar{y}$) = 68

试问我们可否下结论新方法教学效果较好？依假说检定之法则：

1. 我们若以 H_1 表示新方法使用后全国各班的平均成绩，而 H_2 为旧方法的平均成绩，则

$$H_0: H_1 \leq H_2, \text{ 而 } H_a: H_1 > H_2.$$

2. 我们的决策方法自然是若 $\bar{x}$ 比 $\bar{y}$ 大得多，则取 H_a ，现在 $\bar{x}$

比 $\bar{y}$ 大2, 算不算多呢?

3. 我们若在目前的情形下取 H_0 , 则错误的或然率是

$$\alpha = p_r\{\bar{x} - \bar{y} \geq 2 | H_0\} \quad (7)$$

上式是指当 H_0 成立时, 我们若各取4班, 他们会有多少机会产生大于2的平均差。

要计算(7)的或然率相当不容易, 甚至无从下手⁵。但若从模拟实验的观点来看, 因在 $H_1 = H_2$ 时 α 最大, 故

$$\alpha = p_r\{\bar{x} - \bar{y} \geq 2 | H_0\} \leq p_r\{\bar{x} - \bar{y} \geq 2 | H_1 = H_2\} \quad (8)$$

又因在 $H_1 = H_2$ 时, 成绩已与教学法无关, 故这八个数据其实是从同样的教法得出来的。 x_1, x_2, x_3, x_4 可以是其中任何四个数。因此(8)式所要求的是, 若我们从这八个数中任取4个做 x_1, x_2, x_3, x_4 , 而另四个做 y_1, y_2, y_3, y_4 , 有多少可能 $\bar{x} - \bar{y} \geq 2$? 这样的实验很容易做, 经不断反复抽样可得 $\alpha \approx 0.36$ 。因此, 我们若以目前的数据来宣布新方法较好, 可以有高达0.36错误的或然率, 似嫌草率。

§ 3 结 论

由于计算机的进展, 许多困难的统计问题都可以以模拟抽样的方式得到解决。

注1

当我们作统计性的结论的时候, 要特别小心我们样本中所含的结论是不是我们文字所表示的意思。像现在所谈的功课好的小孩子电视看得多, 可以有以下几种结论:

1. 电视看得多的小孩子功课好。
2. 功课好的孩子有多的时间看电视。

3. 允许小孩子多看电视的家庭出功课好的孩子。

而这三种结果对一般人所产生的效果可能不同：第一种结论好像是说看电视会使得功课好；第二种结论就完全没有这种意味；第三种结论是指家庭的教育方式或富裕的情形是功课好与看电视的公因子。而上面结论何者为正确，不是从以上四十个样本只调查功课好坏及看电视多少可以决定的。因此当我们抽完样本之后的正确结论只可能是：

看电视多少与功课好有正的相关性

以上的结论正确的或然率不小于某个百分比。上结论中的“正相关”的意思是指看电视与成绩高低有同升同降的关系。

注2 H_0 又命名为对立假设 (alternative hypothesis) 而 H_0 为虚无假设 (null hypothesis)。

注3 也可以称之为 p -value。在近代统计学中 α 与 p -value 含义不完全相同，但因牵涉太广，不能在此申论。

注4 对这个问题的合理决策可以用好几种，可以是文中的 Q ，但也可以用 $\sum_{i=1}^n (o_i - n/6)^2$ (平方差) 或 $\max(o_i) - \min(o_i)$ (绝对差)，但用模拟法求 α 皆极相似 (但解析法则相差极大，甚至不知如何去解)。

注5 传统的方法是假定 x_i 与 y_i 的分布是常态分布，然后推导出或然率，即所谓的 t -test，或只用大小的次序，即所谓的 nonparametric test。

附注

本文中的一部分系取材于本文著者所著《现代应用数学》一书。